

ASPECTS OF AGENCY

*Decisions, Abilities, Explanations,
and Free Will*



ALFRED R. MELE

Aspects of Agency

Aspects of Agency

Decisions, Abilities,
Explanations, and Free Will

ALFRED R. MELE

OXFORD
UNIVERSITY PRESS

OXFORD

UNIVERSITY PRESS

Oxford University Press is a department of the University of Oxford. It furthers the University's objective of excellence in research, scholarship, and education by publishing worldwide. Oxford is a registered trade mark of Oxford University Press in the UK and certain other countries.

Published in the United States of America by Oxford University Press
198 Madison Avenue, New York, NY 10016, United States of America.

© Oxford University Press 2017

All rights reserved. No part of this publication may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, without the prior permission in writing of Oxford University Press, or as expressly permitted by law, by license, or under terms agreed with the appropriate reproduction rights organization. Inquiries concerning reproduction outside the scope of the above should be sent to the Rights Department, Oxford University Press, at the address above.

You must not circulate this work in any other form
and you must impose this same condition on any acquirer.

Library of Congress Cataloging-in-Publication Data

Names: Mele, Alfred R., 1951– author.

Title: Aspects of agency : decisions, abilities, explanations, and free will / Alfred R. Mele.

Description: New York : Oxford University Press, 2017. |

Includes bibliographical references and index.

Identifiers: LCCN 2016038292 (print) | LCCN 2016039591 (ebook) |

ISBN 9780190659974 (cloth : alk. paper) | ISBN 9780190660000 (online course) |

ISBN 9780190659981 (pdf) | ISBN 9780190659998 (ebook)

Subjects: LCSH: Agent (Philosophy) | Act (Philosophy) | Free will and determinism. |

Ethics. Classification: LCC B105.A35 M438 2017 (print) | LCC B105.A35 (ebook) |

DDC 128/.4—dc23

LC record available at <https://lcn.loc.gov/2016038292>

9 8 7 6 5 4 3 2 1

Printed by Sheridan Books, Inc., United States of America

To my nieces and nephews

Contents

<i>Preface</i>	ix
1. Introduction	i
2. Deciding to Act	7
3. Actions, Explanations, and Causes	27
4. Agents' Abilities	63
5. Free Will and Moral Responsibility: Does Either Require the Other?	91
6. Is What You Decide Ever up to You?	109
7. Arbitrary Decisions and the Problem of Present Luck	135
8. Complete Control and Disappearing Agents	153
9. Libertarianism and Human Agency	181
10. Two Libertarian Theories: Or Why Event-Causal Libertarians Should Prefer My Daring Libertarian View to Robert Kane's View	197
11. Living without Agent Causation	221
<i>References</i>	237
<i>Index</i>	245

Preface

THE ORIGINAL PLAN for this book was to focus on a small collection of important topics in the philosophy of action, basing the chapters on published articles of mine. My original list of topics had seven items on it, the last being free will. My claim was that a good understanding of my six topics would help us grapple with issues that receive much more attention—for example, free will. I thought it would be a good idea to support that claim by including a chapter based on an article of mine on free will that makes use of my work on some of the less-discussed topics. But things quickly got out of hand! After preparing a version of “Is What You Decide Ever up to You?” (Mele 2013b) as a chapter, I thought I should say more about libertarianism (based on Mele 2013c). Once I drafted that chapter, it seemed appropriate to include a version of an article of mine (Mele n.d.c) on a certain argument against a libertarian view that I find more attractive than competing libertarian views. And so on. In the end, I whittled my original list of seven topics down to four: decisions to act (or what I call practical decisions), agents’ abilities, commitments of a causal theory of action explanation, and free will. Seven of this book’s eleven chapters are on free will. Much of this book is based on published articles of mine. References to the article or articles on which a chapter is based appear in a note at the end of the chapter.

This book could easily have grown much longer. An unusual feature of my stand on free will is agnosticism about compatibilism (Mele 1995a, 2006). I have developed a libertarian position on free will for incompatibilists and a compatibilist position for compatibilists while remaining neutral on whether compatibilism or incompatibilism is true. Largely in response to critics, I have published quite a bit on both positions since they were advanced in my *Free Will and Luck* (Mele 2006). Here, to keep this book relatively short, I have decided to focus the discussion of free will on libertarianism. A separate book

on compatibilism is in progress. (The technical terminology in this preface will be explained in chapter 1 for the uninitiated.)

A draft of this book was among the assigned readings for a graduate seminar Michael Robinson and I taught at Florida State University in the fall of 2015. I am grateful to Michael and the students for useful feedback. Thanks are also owed to Ish Haji and Derk Pereboom, both of whom kindly agreed to referee this book. Acknowledgments of other helpful feedback appear in the final endnote of each chapter. This book was made possible through the support of a grant from the John Templeton Foundation. The opinions expressed here are my own and do not necessarily reflect the views of the John Templeton Foundation.

Introduction

I CLOSED MY *Free Will and Luck* (Mele 2006) with the following pair of sentences: “As I write this concluding chapter, ten years have passed since I wrote *Autonomous Agents* [Mele 1995a]. If I am lucky enough to be alive and well ten years from now, perhaps I will take another shot at free will” (p. 207). As readers of my preface know, my original plan for this book was not that it be another shot at free will. It has turned out to be largely that, but I have decided to include detailed discussions of my other three topics rather than saying just enough about them for the purposes of my discussion of free will. Those three topics are decisions to act, commitments of a causal theory of action explanation, and agents’ abilities.

1.1. A Little Background

As I reread the articles of mine on which much of this book is based, I made notes about relatively basic ideas that came up in more than one article. In some cases, I judged it best to introduce the idea in this introductory chapter.

I start with a very brief review of some theories about how actions are individuated. Confusion on this topic can affect judgments about, for example, what an agent is doing for a reason, is doing freely, and is morally responsible for doing. Consider the following from Donald Davidson: “I flip the switch, turn on the light, and illuminate the room. Unbeknownst to me I also alert a prowler to the fact that I am home” (1980, p. 4). How many actions has the agent, Don, performed? Davidson’s *coarse-grained* answer is one action “of which four descriptions have been given” (p. 4; see Anscombe 1963). A *fine-grained* alternative treats *A* and *B* as different actions if, in performing them, the agent exemplifies different act-properties (Goldman 1970). According to this view, Don performs at least four actions, since the act-properties at issue

are distinct. For example, the property of flipping a switch is distinct from the property of turning on a light, and the property of turning on a light (in a room) is distinct from the property of illuminating a room. One may flip a switch without turning on a light, and vice versa. Similarly, one may turn on a light in a room without illuminating the room (the light may be painted black) and illuminate a room without turning on a light (by setting a dark room on fire). Another alternative—a componential one—represents Don's illuminating the room as an action having various components, including (but not limited to) his moving his arm (an action), his flipping the switch (an action), and the light's going on (Ginet 1990; Thalberg 1977; Thomson 1977). Where proponents of the coarse-grained and fine-grained theories find, respectively, a single action under different descriptions and a collection of intimately related actions, advocates of the various componential views locate a "larger" action having "smaller" actions among its parts.

When actions are individuated in Davidson's coarse-grained way, it is *actions under descriptions* that are performed intentionally or unintentionally—not actions period. Similarly, it is actions under descriptions that are performed for a reason *R* (or for a reason at all). For example, under the description "flips the switch," what Don does is intentional; but under the description "alerts the prowler," what he does is not intentional. And under the description "flips the switch," Don might have acted for a reason having to do with getting sufficient light for reading; but under the description "alerts the prowler," Don does not act for a reason at all (Davidson 1980, p. 5).

Fine-grained and componential theorists make no special appeal to action-descriptions. Theorists of both kinds can straightforwardly say, for example, that Don intentionally flipped the switch and unintentionally (or nonintentionally) alerted the prowler.¹ They can also say that Don flipped the switch in order to have enough light for reading and that, although he alerted the prowler, he did not do that for a reason.

The nature of action is next. Philosophers of action try to avoid using the word "action" indiscriminately. In ordinary English, people speak not only of the actions of human beings and other intelligent animals but also of the actions of acids, winds, and waves. Acids dissolve things, winds blow things around, and waves push and drag things; and these events count as actions in a broad sense. Mainstream philosophy of action is not concerned with actions of inanimate objects, and its primary subject matter is *intentional* action. As Davidson understands action, *every* action is intentional under some acceptable description (also see Hornsby 1980). (On this view, if an event is not intentional under any acceptable description, that event is not an action.)

Davidson contends that “a man is the agent of an act if what he does can be described under an aspect that makes it intentional” (1980, p. 46) and that “action . . . require[s] that what the agent does is intentional under some description” (p. 50). Putting these remarks together, we get the thesis (*TD*) that *x* is an action if and only if *x* is intentional under some description.²

Proponents of alternative theories of action individuation may embrace the following rough analogue of *TD*: (*TA*) In every case of action something is done intentionally; when nothing is done intentionally, no action is performed. Notice that in the story in which Don unknowingly alerts the prowler, something relevant is done intentionally: for example, Don intentionally flips the switch. If Don were sound asleep and were to alert the prowler by snoring, his alerting the prowler would not be an action—not even an unintentional (or nonintentional) one, according to this analogue of Davidson’s view. (On Davidson’s own view, because there is no acceptable description of what Don is doing while he is sleeping under which it is intentional, “Don’s alerting the prowler” does not describe an action.)

I am neutral regarding the three theories of action individuation I sketched. Henceforth, readers should understand the action variable *A* as a variable for actions themselves (construed componentially or otherwise) or actions under descriptions, depending on their preferred theory of action individuation. The same goes for the expressions that take the place of *A* in concrete examples. For example, fans of Davidson’s theory should read “Don flips the switch” as “something Don does under the description ‘flips the switch’” and non-Davidsonians should make no adjustments.

Intentions receive considerable attention in this book. It is useful to distinguish among proximal, distal, and mixed intentions (Mele 1992a, pp. 143–44, 158). Angela has an intention to phone Nick right now. This is a *proximal* intention. Proximal intentions also include intentions to continue doing something that one is doing and intentions to start *A*-ing (for example, start running a mile) straightaway. *Distal* intentions are intentions for the non-immediate future—for example, Al’s intention to shoot pool with Jack tomorrow night and Ann’s intention to phone Dave at the predesignated time, exactly one minute from now. Some intentions have both proximal and distal aspects. For example, Ann may have an intention to run a mile without stopping, beginning now. (She estimates that the deed will take five minutes.) I call such an intention a *mixed* intention. An intention of this kind specifies something to be done now and something to be done later. Exactly parallel distinctions may be made in the case of desires and decisions. There are proximal, distal, and mixed desires and decisions.

I close this section with some brief remarks on the expression “free will” and on some common terminology in the literature on free will. I conceive of free will as the ability to act freely and treat free action as the more basic notion. But what is it to act freely? As I observe in Mele 2006, there are readings of “freely *A*-ed” on which the following sentence is true: “While Bob was away on vacation, mice ran freely about his house.” Such readings do not concern me. My interest is in what I call *moral-responsibility-level free action*—“roughly, free action of such a kind that if all the freedom-independent conditions for moral responsibility for a particular action were satisfied without that sufficing for the agent’s being morally responsible for it, the addition of the action’s being free to this set of conditions would entail that he is morally responsible for it” (Mele 2006, p. 17).³

I have mentioned compatibilism, incompatibilism, and libertarianism. Determinism enters into all three ideas. Peter van Inwagen describes determinism as “the thesis that there is at any instant exactly one physically possible future” (1983, p. 3). The thesis he has in mind, expressed more fully, is that at any instant exactly one future is compatible with the state of the universe at that instant and the laws of nature. There are more detailed characterizations of determinism in the literature; but this one is fine for my purposes. (An exception may be made for instants at or very near the time of the Big Bang.)

Compatibilism (about free will) is the thesis that free will is compatible with determinism. In terms of possible worlds, compatibilism is the thesis that there are possible deterministic worlds in which free will exists. Incompatibilism is the denial of compatibilism.⁴ Libertarianism is the conjunction of two theses:

1. *The incompatibility thesis*: Free will is incompatible with determinism.
2. *The pro-free-will thesis*: There are actions that are or involve exercises of free will—*free actions*, for short.

I have much more to say about libertarianism than compatibilism in this book. Different libertarians take different positions on what free will requires beyond the falsity of determinism. *Agent-causal* libertarians contend that only beings with agent-causal powers can have free will. *Noncausal* libertarians argue that only uncaused actions can be (directly) free.⁵ *Event-causal* libertarians avoid appealing to agent causation, and they typically claim that paradigmatically free actions are indeterministically caused by their proximal causes. For reasons that will emerge, I regard this third brand of libertarianism

as by far the most promising brand, and I will defend its positive side against a variety of objections.⁶

Among the most dogged critics of event-causal libertarianism are philosophers who contend that possessing the power of agent causation is required for having free will (see Clarke 2003; O'Connor 2000; Pereboom 2001, 2014). Agent causation is causation of an effect by an agent or person, as opposed to causation of an effect by states or events of any kind, including a person's motivational and representational states. Agent causation is not reducible to causation by events or states. Most agent causationists prefer their agent causation straight (Chisholm 1966; O'Connor 2000; Taylor 1966), but it may be mixed with event causation in a theory about the production of free actions (Clarke 2003).

1.2. Preview

I close this chapter with a brief preview of the remainder of this book. Chapter 2 defends an account of what it is to decide to do something and makes a case for the claim that there are genuine actions of decision-making. The topic of chapter 3 is a project that is not as well understood as it should be—the project of constructing a causal theory of the explanation of intentional actions. My aim there is to clarify the commitments of such a theory. Chapter 4 distinguishes among three different kinds or levels of agents' abilities.

In chapter 5, I turn to free will. That chapter explores the conceptual connections between free action and action for which an agent is morally responsible. Chapter 6 takes up the question what it might be for it to be *up to* an agent what he does. The chapter highlights and defends a plank in a libertarian position I floated in *Free Will and Luck* (Mele 2006) in response to a certain worry about luck. The main business of chapter 7 is a critique of some familiar control-featuring arguments against event-causal libertarianism. In chapter 8, I explain why readers should not be persuaded by Derk Pereboom's "disappearing agent objection" to event-causal libertarianism (2014, chap. 2) while also exploring a notion of complete control over whether one will decide to *A*. Chapter 9 investigates the contribution indeterministic agent-internal processes of a kind for which we have some indirect evidence might make to free will beyond being sufficient for the falsity of determinism. The thesis of chapter 10 is that event-causal libertarians should prefer the "daring libertarian" view that I floated in *Free Will and Luck* to Robert Kane's well-known

concurrent-efforts event-causal libertarian view (1999b, 2011, 2014). Chapter 11 closes the book with a speculative discussion of how philosophers might proceed if it were to be proved that agent causation is conceptually impossible.

Notes

1. On the existence of a middle ground between intentional and unintentional actions, see Mele 2012b. Actions on this middle ground may be called *nonintentional*.
2. Davidson expresses the point differently: “a person is the agent of an event if and only if there is a description of what he did that makes true a sentence that says he did it intentionally” (1980, p. 46).
3. The subjunctive conditional I quoted leaves it open that there are moral-responsibility-level free actions for which the agents are not morally responsible. As I understand moral responsibility, agents are not morally responsible for nonmoral actions (see chap. 5, sec. 5.2); and a nonmoral action may satisfy the subjunctive conditional.
4. For nontraditional uses of “compatibilism” and “incompatibilism,” see Mickelson 2015.
5. It is open to a noncausal libertarian to claim that some caused actions are indirectly free and inherit their freedom from the freedom of some uncaused free actions that are among their causes (see chap. 5, sec. 5.1).
6. For overviews of the first two kinds of view, see, respectively, O’Connor 2011 and Pink 2011.

Deciding to Act

Each of us makes decisions all day long. We choose which clothes to wear in the morning. We pick a route to take to work. We plan how much money to withdraw from the bank and set aside for food, entertainment, and incidentals. On some days, we even make “monumental” decisions, for example, what job to take, what car to buy, whether or not to get married, and to whom.

—(Yates 1990, pp. 2–3)

AS THIS PASSAGE from a book on the psychology of decision-making indicates, deciding seems to be part of our daily lives. But what is it to decide to do something? It may be true, as some philosophers have claimed, that to decide to *A* is to perform a mental action of a certain kind—specifically, an action of *forming* an intention to *A*.¹ (The verb “form” in this context is henceforth to be understood as an action verb.) Even if this is so, there are pressing questions. Do we form all of our intentions? If not, how does forming an intention differ from other ways of acquiring one? Do we ever form intentions, or do we rather merely acquire them in something like the way we acquire beliefs or desires? These are among the focal questions of this chapter. My aim is to clarify the nature of deciding to act and to make a case for the occurrence of genuine actions of intention formation.

2.1. Background: Four Views of Practical Deciding

We speak not only of deciding to act but also of deciding that something is the case, as an economist, on the basis of careful research, might decide that oil prices are likely to rise dramatically in the next few months. For stylistic purposes, it is useful to have labels for the two kinds of deciding. Following Arnold Kaufman (1966, p. 25), I call deciding to act *practical* deciding and deciding that something is the case *cognitive* deciding. I also count decisions *not* to *A*—for example, not to vote in an upcoming election—as practical decisions. My focus, however, is deciding to act.

Are instances of practical deciding *actions*, or is action the wrong category for them? A brief sketch of a familiar nonactional conception of cognitive deciding will help to explain why someone may take practical deciding to be nonactional. On one possible view, to decide that p is the case is simply to acquire a belief that p is the case on the basis of reflection. For example, the economist's deciding that oil prices are likely to rise soon is a matter of his acquiring the belief that this is likely on the basis of reflection on considerations that he takes to be relevant. The economist's belief is a product of his reflective activity, but it is not as though his acquiring that belief is itself an action. It is not as though, on the basis of his reflection, he performs an action of belief formation. In some spheres, acquiring an x may have both an actional and a nonactional mode. The economist's acquiring a car, for example, may or may not be an action of his. In buying my car, he performs an action of car acquisition; if, instead, I give him my car, his acquiring it is not an action of his. According to a nonactional view of cognitive deciding, acquiring a belief on the basis of reflection is never an action.

View 1: practical deciding as nonactional. An analogous view of *practical* deciding is possible. According to this view, to decide to A is simply to acquire an intention to A on the basis of practical reflection, and acquiring an intention—in this way or any other—is never an action. To be sure, many intentions are products of reflective activity, but that is not to say that there are any actions of intention acquisition. Buying a car is an action of car acquisition, as is stealing a car. But there are no analogues of these activities in the realm of intention acquisition.²

View 2: practical deciding as extended action. A variant of this view that represents practical deciding as actional is imaginable. The view just sketched highlights a process that culminates in intention acquisition. The process centrally involves practical reflection, mental activity. It may be suggested that “practical deciding” is a name for the whole process. In any case in which an agent (nondeviantly) acquires an intention on the basis of practical reflection, there is a process of practical deciding that begins when the reflection begins and ends with the acquisition of the intention. Since much of that process is actional, and since the process is a unified whole, it is claimed that the process itself is properly counted as an action—specifically, an action of deciding to do (or not to do) something. A proponent of view 2 has the option of denying or affirming that practical deciding as extended action is sometimes associated with practical deciding of the kind at issue in views 3 and 4. (For more on this, see note 5.) A pure version of view 2 includes the denial of this proposition. It limits practical deciding to the process described here.

View 3: practical deciding as mythical. A skeptical variant of view 1 also is conceivable. Practical deciding, it may be insisted, is essentially actional, but not in the way claimed by view 2. Deciding to *A*, by definition, is a momentary mental action of intention formation. Practical reflection is not part of any action of deciding, although such reflection may often precede and inform instances of practical deciding. However, there is no such thing as practical deciding. For, as view 1 asserts, no instance of acquiring an intention is an action. In this respect, intention acquisition is like belief acquisition, on a plausible, nonactional conception of the latter.

View 4: practical deciding as a momentary mental action of intention formation. A fourth view is view 3 without the skepticism.

I will comment on each of these views in turn. A proponent of view 1 owes us an account of how intentions are, or might be, acquired on the basis of practical reflection, reflection about what to do. No consensus has been reached about the form such reflection takes. As I understand practical reflection or reasoning, it is, roughly, reflection about what to do that is sustained by motivation to answer a practical question (Mele 1992a, chap. 12; 1995a, chap. 2; 2003a, chap. 4). Reflection about what to do sometimes takes the form of reasoning about what it would be best to do. It takes other forms as well. For example, an agent who is inclined to *A* may reflect about whether to *A*, and that reflection may take the form of reasoning about whether *A*-ing is acceptable or “good enough” (Mele 2003a, chap. 4). And someone who intends to *A* may reflect practically on acceptable means to *A*-ing.

Practical reflection, as I understand it, often issues in evaluative beliefs about action: for example, a belief that it would be best to *A*, a belief that *A*-ing would be acceptable, a belief that *A*-ing is a satisfactory means to a desired end (Mele 1992a, chap. 12; 1995a, chap. 2; 2003a, chap. 4). Such beliefs, based as they are on reflection, are cognitive decisions. One might argue that cognitive decisions of the kinds just mentioned often issue directly in corresponding intentions, without any action of intention formation, and that acquiring an intention on the basis of practical reflection is just a matter of an intention’s being directly produced by the acquisition of a reflective evaluative belief. That, it may be alleged, is how intentions are acquired on the basis of reflection. Accordingly, it may be claimed that instances of practical deciding really are not *actions* of deciding at all; rather, they are the nonactional production of intentions by instances of cognitive deciding that are based on practical reflection. It may be claimed, for example, that if an agent judges that it would be best to *A*, that judgment’s issuing, without any *action* of practical deciding, in an intention to *A* is an instance of deciding to *A*.

There is an element of truth in this. In some cases, having judged or decided on the basis of practical reflection that it would be best to *A*, one seemingly does not need to proceed to do anything to bring it about that one intends to *A*. The judgment may issue straightaway and by default in the intention. So, at least, I have argued elsewhere (Mele 1992a, chap. 12). However, things are not always so simple. Consider Joe, a smoker. On New Year's Eve, he is contemplating kicking the habit. Faced with the practical question what to do about his smoking, Joe is deliberating about what it would be best to do about this. It is clear to him that it would be best to quit smoking at some point, but as yet he is unsure whether it would be best to quit soon. Joe is under a lot of stress, and he worries that quitting smoking might drive him over the edge. Eventually, he decides that it would be best to quit—permanently, of course—by midnight. Joe's cognitive decision settles an evaluative question. But Joe is not yet settled on quitting. He tells his partner, Jill, that it is now clear to him that it would be best to stop smoking, beginning tonight. She asks, "So is that your New Year's resolution?" Joe sincerely replies, "Not yet; the next hurdle is to *decide* to quit. If I can do that, I'll have a decent chance of kicking the habit."

This little story is coherent. In some instances of akratic action, one intends to act as one judges best and then backslides (Mele 1987, 2012a). In others, one does not progress from judging something best to intending to do it.³ Seemingly, having decided that it would be best to quit smoking, Joe may or may not form the intention to quit. His *forming* the intention, as opposed to his nonactionally acquiring it, apparently would be a momentary mental action of the sort highlighted in views 3 and 4. (Recall my convention of using the verb "form" in "form an intention" as an *action* verb.) Whether there are such actions remains to be seen.

I turn to view 2, "practical deciding as extended action." It is best construed as a reformative view, as opposed to an explication of a common-sense notion of deciding. A student who says that he "was up all night deciding to major in English" is speaking loosely, but we understand what he means. I would put it this way: he was up all night deliberating about what major to declare, or about whether to major in English, and eventually decided (or acquired an intention) to declare English as a major on the basis of his deliberation.⁴ However, on view 2, claims of the sort this student made are often true on a *literal* interpretation: according to this view, an instance of deciding, an extended action, begins when relevant practical reflection begins and ends when the agent (nondeviantly) acquires an intention on the basis of that reflection.

If there are events that conform to the common-sense notion of deciding at issue in views 3 and 4, view 2 is not particularly interesting. Where there

is no need for conceptual reform, there should be little interest in pursuing it. Whether further discussion of view 2 is in order will depend on what the ensuing discussion of views 3 and 4 uncovers.⁵

According to view 3, practical deciding is a momentary mental action of intention formation and there are no such actions. Return to Joe. Although he decided that it would be best to quit smoking, he does not yet intend to quit. Can he proceed to *form* an intention to quit? Conventional wisdom suggests that he can. But it is difficult to be confident that conventional wisdom is correct about this independently of a good grasp of what it is to form an intention. Intentions themselves are relatively well understood (see Brand 1984; Bratman 1987; Mele 1992a); in section 2.3, I sketch my own understanding of their constitution. The question now is what it is to *form* an intention as opposed to nonactionally acquiring one.

Does one want to understand the *means* by which agents form intentions? No. Assume that there are *basic actions*—roughly, actions that an agent performs, but not by means of doing something else. If there are momentary actions of intention formation, they fall into this group. Perhaps the best way to get a handle on what forming an intention might be is to catalogue ways in which intentions arguably are nonactionally acquired and to see what conceptual space might remain for the actional acquisition of intentions.

I have already commented on one way in which an agent arguably may acquire an intention without forming it: perhaps a cognitive decision of a certain kind may issue directly in a corresponding intention. There may be other ways, as well. Consider the following claim: “When I intentionally unlocked my office door this morning, I intended to unlock it. But since I am in the habit of unlocking my door in the morning, and conditions this morning were normal, nothing called for a *decision* to unlock it” (Mele 1992a, p. 231). It seems that some intentions arise as part of a routine without having to be actively formed. If I had heard a terrible ruckus in my office, I might have paused to consider whether to unlock the door or call the campus police, and I might have decided to unlock the door. But given the routine nature of my conduct, I see no need to postulate an action of intention formation in this case. Moreover, if an action of intention formation must occur in a simple case of this kind to account for my action of unlocking the door, why not suppose that the former action must be the product of another one—a decision to form the intention to unlock the door? An infinite regress threatens.

Arguably, some intentions may nonactionally arise out of desires (Audi 1993, p. 64). Suppose that an agent acquires a proximal desire to *A*. Perhaps, if the agent has no (significant) competing desires and no reservations about *A*-ing,

the acquisition of that desire may straightaway give rise to the nonactional acquisition of a proximal intention to *A*. Walking home from work, Helen notices her favorite brand of beer on display in a store window. The sight of the beer prompts a desire to buy some, and her acquiring that desire issues straightaway in an intention to buy some. This seems conceivable.

It also seems conceivable that, given Helen's psychological profile, the sight of the beer in the window issues *directly* in an intention to buy some, in which case there is no intervening desire to buy the beer (see Mele 2003a, pp. 171–72). Perhaps in some emergency situations, too, a perceptual event, given the agent's psychological profile, straightaway prompts an intention to *A*. Seeing a dog dart into the path of his car, an experienced driver who is attending to traffic conditions may immediately acquire an intention to swerve. This seems conceivable too (see Mele 2003a, p. 185).

Brian O'Shaughnessy, a proponent of view 1 ("practical deciding as non-actional"), claims that decidings are "those comings-to-intend events that resolve a state of uncertainty over what to do" (1980, vol. 2, p. 297). This claim may be divorced from his commitment to view 1. If the basic point about resolving practical uncertainty is correct, it may be correct even if practical decidings are momentary mental actions. Notice also that if the point is correct, it helps to account for common intuitions about scenarios of the kind discussed in the preceding three paragraphs.⁶ In the cases I described, there is no uncertainty that intention acquisition resolves. I was not uncertain about whether to unlock my door, Helen was not uncertain about whether to buy the beer, and the driver was not uncertain about what course of action to take. At no point in time were any of us uncertain about the matters at issue. Furthermore, if there are cases in which a cognitive decision based on practical reflection issues directly (and therefore without the assistance of an act of intention formation) in a corresponding intention, the agent's reaching his cognitive conclusion resolves his uncertainty about what to do. Reaching the conclusion directly results in settledness on a course of action (or, sometimes, in settledness on *not* doing something). In Joe's case, of course, matters are different: even though he has decided that it would be best to quit smoking, he continues to be uncertain or unsettled about what to do.⁷

Having some sense of various ways in which intentions might arise without being formed, why should one think that we sometimes form intentions? As a proponent of view 3 or 4 might put the question, why should one think that we *decide* to do (or not to do) things?

It is natural to consider ordinary experiences of agency in this connection. Many people say they have robust experiences of deciding to act, of making

up their minds to do things. A proponent of view 1 or 3 may argue that what these people in fact experience is cognitive deciding about action: for example, arriving at the belief that it would be best to *A* on the basis of practical reflection. It may also be suggested that people sometimes experience the sort of intention acquisition that occurs in the beer scenario or in some emergency situations and mistakenly take themselves to have *formed* those intentions in an action of practical decision-making. However, the experience claims cannot safely be quickly dismissed.

I will report on some common experiences of mine and then try to ascertain whether they might be veridical. Sometimes I find myself with an odd hour or less at the office between scheduled tasks or at the end of the day. Typically, I briefly reflect on what to do then. I find that I do not try to ascertain what it would be *best* to do at those times: this is fortunate, because settling that issue might often take much more time than it is worth. Instead, I look at a list that I keep on my desk of short tasks that need to be performed sooner or later—reply to an e-mail message, write a letter of recommendation, and the like—and decide which to do. So, at least, it seems to me. Sometimes I have the experience not only of settling on a specific task or two but also, in the case of two or more tasks, of settling on a particular order of execution.

I have an e-mail system that makes a sound when a message arrives. Occasionally, when I hear that sound, I pause briefly to consider whether to stop what I am doing and check the message. Sometimes I have the experience of deciding to check it, or the experience of deciding not to check it. Sometimes I do not even consider checking the new message.

In situations of both of the kinds under consideration (the odd hour and incoming e-mail), I sometimes have the experience of having an urge to do one thing but deciding to do another instead. For example, when I hear that a new e-mail message has arrived, I may have an urge to check it straightaway but decide to finish what I am doing first. (When I am grading papers, these urges tend to be particularly strong.) When I am looking at my list of short tasks at the beginning of an odd hour, I may feel more inclined to perform one of the more pleasant tasks on my list but opt for a less pleasant one that is more pressing.

That I lead an exciting life is obvious. The question now is whether my reported experiences ever match reality. Sometimes, it seems, practical uncertainty is resolved by our arriving, on the basis of practical reflection, at cognitive decisions that issue directly in corresponding intentions. Do we also resolve practical uncertainty by deciding—in the momentary mental action sense—what to do? If my utterly mundane experiences in my office can be trusted,

the answer is *yes*. But can they? A persuasive response to this last question requires a careful look at some sources of skepticism about the existence of practical decisions, construed as momentary mental actions.

2.2. Intentions in Practical Deciding

In this section, I assume that practical decidings are momentary mental actions of intention formation and investigate a problem for the thesis that practical deciding, so understood, is a genuine phenomenon. That there are mental actions should be uncontroversial. If you feel inclined to demur, please take a break from this chapter and solve the following multiplication problem in your head: $157 \times 15 = ?$ Unless you are uncommonly talented in this sphere, it will take you at least a few seconds to succeed. If you persist, you will perform the mental action of solving the problem in your head; and even if you give up too soon, you will have performed mental actions in trying to solve the problem in your head. That practical deciding is a mental action is not a mark against it. But some other aspects of it might be.

One potential source of skepticism about the existence of practical decisions is a worry about their etiology. To decide to *A*, according to the working assumption of this section, is to form an intention to *A*. Does deciding to *A* require an intention in addition to the intention to *A* formed in so deciding? As I mentioned in chapter 1, Donald Davidson has claimed that every action is intentional under some description or other (1980, chap. 3). According to another popular thesis, in every case of overt intentional action, some intention or other plays an action-producing role (Brand 1984, chap. 2; Bratman 1987, pp. 119–27; Mele 1992a, chap. 10). Assuming that practical decidings are actions, Davidson's thesis entails that they are intentional under some description. In that case, a counterpart of the second thesis would imply that in every case of deciding to *A*, some intention or other plays a productive role. If it is plausible that our best account of overt intentional action places intentions in causal roles and that practical decisions are mental analogues of some overt intentional actions, we should expect intentions to play a role in the production of practical decisions. If it should turn out that there are no plausible candidates for intentions that play such a role in garden-variety practical decision-making, there may be grounds for skepticism about the existence of practical decisions, construed as actions.

Hugh McCann has defended the view that practical decisions are intentional actions (1986a). He also contends that to say that an agent's decision to *A* "is intentional cannot mean that he made it out of a prior intention so

to decide, for if he already intended to decide that way then he had already decided" (1986a, p. 266).⁸ Although this is not quite right, McCann is on to something. Seeing where he goes wrong will prove useful.

Can someone intend to decide to *A* without already having decided to *A*? Consider the following story. Brian is deliberating about whether to *A* or to *B*. A demon has manipulated Brian's brain in such a way that he temporarily cannot do any of the following: decide to *A*, decide to *B*, and nonactionally acquire either intention (that is, the intention to *A* or the intention to *B*). The demon informs Brian that if he is to decide to *A* he must first press a certain green button that will enable him to decide to *A* in just the way he decides to do other things and that pressing a certain blue button will enable him to decide to *B* in the normal way. Brian continues deliberating and eventually comes to the conclusion that it would be best to *A*. In other circumstances, his so judging might have issued straightaway in an intention to *A*, but the demon has prevented that from happening. Judging it best to *A*, and believing that it is very unlikely that he will *A* without intending to, Brian wants to form the intention—that is, to decide—to *A*. In fact, he intends to decide to *A*. Believing, correctly, that he must press the button in order to enable himself to decide to *A*, he presses the button. And then he decides to *A*.

Is this story coherent? Is there a hidden contradiction, or perhaps a contradiction that is evident to everyone but me? One might claim (1) that intending to decide to *A* is conceptually sufficient for being settled on *A*-ing and (2) that being settled on *A*-ing is conceptually sufficient for intending to *A*. If claims 1 and 2 are true, the story has a large hole in it. According to the story, Brian cannot intend to *A* until he presses the green button. But claims 1 and 2 and the detail that Brian intends to decide to *A* before he presses the button jointly entail that Brian has the intention to *A* before he presses the button.

The culprit here is claim 1, not the story. Claim 1 is part of what the story is designed to test. Grant that the demon's machinations will prevent Brian from being settled on *A*-ing unless and until he presses the green button. Does it follow from this that Brian cannot intend to decide to *A*? Does it follow that he cannot intend to bring it about that he is settled on *A*-ing? I do not see how. Being settled on *A*-ing is one thing, and being settled on bringing it about that one is settled on *A*-ing is another. If being in the latter, higher-order condition were to entail being in the former, lower-order condition, only confused agents could be in the higher-order condition. Plainly, an agent who is settled on bringing it about that he is settled on *A*-ing is under the impression that he is not yet settled on *A*-ing. It is difficult to see why it should be thought that his impression *must* be mistaken. McCann urges,

reasonably, that a psychological commitment to *A*-ing is central to intending to *A* (1986a, p. 251).⁹ That the demon has set up an obstacle to Brian's being psychologically committed to *A*-ing is compatible with Brian's being psychologically committed to removing the obstacle and to bringing it about that he is psychologically committed to *A*-ing.

Although this is, to be sure, a highly contrived case, it undermines McCann's claim. More important, it helps set the stage for a useful positive observation. The idea that decisions to *A* normally are produced in part by intentions to decide to *A* is indeed odd and unacceptable, as I will explain shortly. Another bit of stage-setting is in order first.

If an agent acquires an intention to decide to *A*, either he acquires it nonactionally or he decides to decide to *A*. The latter disjunct points the way to an impossible regress that can be cut short by supposing that at some point the agent nonactionally acquires a pertinent "decisional" intention—say, the intention to decide to decide to *A*. Since we wind up with a nonactionally acquired decisional intention in either case, we do better to consider the former disjunct.

Now, why might it happen that an agent nonactionally acquires an intention to *decide* to *A* rather than nonactionally acquiring an intention to *A*? In a normal scenario, in which there is no special payoff specifically for deciding to *A* or for intending so to decide, an agent's nonactionally acquiring an intention to decide to *A* when he could just as easily have nonactionally acquired an intention to *A* would be notably inefficient.¹⁰ For example, having acquired a *proximal* intention to *A* (in this case, an intention to *A* right then), he would proceed to *A* straightaway, if all goes smoothly; whereas having acquired a proximal intention to *decide* to *A* straightaway, he would, if all goes smoothly, make the decision before he *A*-s. Similarly, acquiring at *t* a distal intention to *A* is an instance of becoming settled at *t* on *A*-ing later, whereas acquiring at *t* an intention to decide (right then or later) to *A* later is the sort of thing that would lead to settledness on *A*-ing by way of an additional step—decision-making. Setting aside the above-mentioned "special payoff" scenarios, I conjecture that an agent would nonactionally acquire an intention to decide to *A* only if there were some obstacle to his nonactionally acquiring an intention to *A* that is not also an obstacle (or not as great an obstacle, or not accompanied by as great an obstacle) to his nonactionally acquiring an intention to decide to *A*. There is such an obstacle in my demon story. Other scenarios featuring an obstacle of the kind at issue would, I believe, also be highly unusual. In normal cases, there is no purpose in nonactionally acquiring an intention to decide to *A* that is not more efficiently served by a nonactional acquisition of

an intention to *A*. In the absence of a reason to suppose that our nonactional mental life normally is inefficiently complicated in this way, the assumption of greater efficiency and simplicity is very plausible.

If intentions to decide to *A* normally do not play a role in producing our decisions to *A*, should we abandon the idea that we perform actions of practical decision-making? Might intentions of another kind be useful in this connection? Arguably, there are a great many things that we do intentionally without having intentions specifically to do them. When I walk to work, for example, my individual steps are intentional actions, but there is no need to suppose that each step requires its own distinct intention. In a typical case, a single, more general intention—an intention to walk to work along my normal route—is intention enough. Similarly, when a man intentionally runs a mile, he runs its various parts (for example, the first quarter-mile, the seventh eighty-yard segment), and his runnings of those segments are intentional actions; but there is no need to suppose that he has distinct intentions for each of these segments.¹¹

If one can *A* intentionally without having an intention specifically to *A*, perhaps one can decide to *B* without having an intention to decide to *B*. If a certain kind of intention is normally at work in the production of decisions, it might be an intention to decide what to do (Mele 1997, p. 243; see Kane 1996, pp. 138–39). In a normal agent who decides to *B* while lacking an intention to decide to *B*, an intention to decide what to do may make an important contribution to the production of that decision. The idea that in every instance of practical deciding some intention or other plays a productive role does not have the unfortunate implication that decisions to *B* always or normally are produced (in part) by intentions to decide to *B*.

This section's topic has been a specific challenge to the view that there are momentary mental actions of practical decision-making. The challenge rests on two assumptions: (*A1*) practical decidings are intentional actions; (*A2*) intentions play a role in the production of all intentional actions. The worry was that there are no plausible candidates for intentions that play a productive role in garden-variety practical decision-making. I endorse *A1*. Why haven't I explicitly endorsed *A2*? The answer is partly pragmatic. Although I am inclined to accept a version of *A2* (partly for reasons of simplicity), a proper defense of it would require a chapter of its own, and I have decided against writing such a chapter. More important, the adequacy of my response to the worry addressed in this section does not depend on the truth of *A2*; rather, the worry itself derives largely from *A2*. If *A2* is false, the worry may be dismissed on the grounds that it rests on a falsehood. So suppose that *A2* is true.

In that event, I have suggested, there is a kind of intention that is plausibly involved in the production of garden-variety practical decisions—an intention to decide what to do. (Since practical decisions include decisions *not* to *A*, the intention may be described more fully as an intention to decide what to do or not to do.)¹²

2.3. Practical Deciding

In an attempt to resolve a particular practical uncertainty—about whether to *A* or *B*, for example—we may do a variety of things. We may gather relevant data, or hire someone to gather data; we may assess data, or pay someone for a professional assessment; and so on. However, even an agent who is convinced both that all the relevant information is in and that further assessment or deliberation would prove fruitless may be unsettled about what to do. Such an agent might consider attempting to resolve his practical uncertainty by deciding to abide by the toss of a coin and then tossing it. But if a roundabout attempt of this kind at resolving uncertainty is open to him, what is to prevent him from making a direct attempt? (It might even occur to the agent that if he can decide to abide by the result of a coin toss, he can, in the same way, decide to *A*, or decide to *B*.) Well, what would the directness of an attempt to resolve uncertainty amount to in a situation of the imagined kind? A philosopher who cannot answer this question is in no position to treat the preceding one as rhetorical.

This brings me back to my original question. What is it to decide to do something? In section 2.2, I argued that a view of kind 4 (“practical deciding as a momentary mental action of intention formation”) survives the most pressing worry about it. In this section, I articulate the view more fully and tie together the main strands of argument in earlier sections. I begin with a thumbnail sketch of intention.

Intentions, on my view, are executive attitudes toward plans (Mele 1992a, chaps. 8–11).¹³ Like beliefs and desires, intentions have representational content. The representational content of an intention is an action-plan. In the limiting case, the plan-component of an intention has a single “node.” It is, for example, a prospective representation of one’s taking a vacation in Lisbon next winter that includes nothing about means to that end nor about specific vacation activities. Often, intention-embedded plans are more complex. The intention to check her e-mail that Jan executed this evening incorporated a plan that included clicking on her e-mail icon, then typing her password in a certain box, then clicking on the “OK” box, and

so on. An agent who successfully executes an intention is guided by the intention-embedded plan.¹⁴

Although the contents of intentions are plans, I follow the standard practice of using such expressions as “Jan’s intention to check her e-mail now” and “Ken intends to bowl tonight.” It should not be inferred from such expressions that the agent’s intention-embedded plan has a single node—for example, checking e-mail now or bowling tonight. Often, our expressions of an agent’s desires and intentions do not identify the full content of the attitude and are not meant to. Jan says, without intending to mislead, “Ken wants to bowl tonight,” knowing full well that what he wants is to bowl with her at MegaLanes tonight for \$20 a game until the place closes, as is their normal practice.

According to a popular view of representational attitudes—for example, Lara’s belief that *p*, Mel’s desire that *p*, Nora’s desire to *A*, Owen’s intention to *A*—one can distinguish between an attitude’s representational content and its psychological *orientation* (Searle 1983). Orientations include believing, desiring, and intending. On my view, the executive dimension of intentions is intrinsic to the attitudinal orientation *intending*. We can have a variety of attitudes toward plans: for example, we might admire plan *x*, be disgusted by plan *y*, and desire to execute plan *z*. To have the intending attitude toward a plan is to be settled (but not necessarily irrevocably) on executing it.¹⁵ The intending and desiring attitudes toward plans differ in that the former alone entails this settledness. Someone who desires to *A*, or to execute a certain plan for *A*-ing—even someone who desires this more strongly than he desires not to *A*, or not to execute that plan—may still be deliberating about whether to *A*, or about whether to execute the plan, in which case he is not settled on *A*-ing, or not settled on executing the plan.¹⁶ Pat wants more strongly to respond in kind to a recent insult than to refrain from doing so, but, owing to moral qualms, she is deliberating about whether to do so. She is unsettled about whether to retaliate despite the relative strength of her desires (see Mele 1992a, chap. 9).

On a standard view of desire, the psychological features of desires to *A* in virtue of which they contribute to intentional *A*-ings are their content and their strength. On my view of the contribution of *intentions* to *A* to intentional *A*-ings, the settledness feature of intentions is crucial, and it is not capturable in terms of desire strength (and content), nor in terms of this plus belief (Mele 1992a, pp. 76–77 and chap. 9). Intentions to *A*, as I understand them, essentially encompass motivation to *A*, but without being reducible to a combination of desire and belief (Mele 1992a, chap. 8). Part of what it is to be

settled on *A*-ing is to have a motivation-encompassing attitude toward *A*-ing; lacking such an attitude, one lacks an element of a psychological commitment to *A*-ing that is intrinsic to being settled on *A*-ing, and therefore to having an intention to *A*.

There is a body of literature on belief constraints on intentions. Elsewhere, I have defended the view that *S* intends to *A* only if *S* lacks the beliefs that he will not *A* and that he probably will not *A* (Mele 1992a, pp. 146–51). If this is right (but, for the purposes of this book, I have no need to insist that it is), an agent may have an executive attitude toward a plan for *A*-ing without having an intention to *A*. For example, a golfer may have an executive attitude toward a plan for sinking a putt that he believes he probably will miss. But this is not to say that the attitude at issue is not an intention. It is plausibly deemed an intention to try to sink the putt.

Assume that intentions to act are executive attitudes toward plans. As I have argued, one can distinguish between *merely acquiring* such an attitude and *actively settling* on a course of action. Given the assumption just made, deciding to *A* is the latter. Deciding to act—actively settling on a course of action—may be understood as a mental action of executive assent to a first-person plan of action. If you tell me that Mike is an excellent basketball player and I express complete agreement, I thereby assent to your claim. This is overt *cognitive* assent. If you propose that we watch Mike play tonight at the arena and I express complete acceptance of your proposal, I thereby assent to your proposal. This is overt *executive* assent: I have agreed to join you in executing your proposal for joint action. Now, perhaps my overt action of assenting to your proposal was a matter of my giving voice to a nonactionally acquired intention to join you in watching Mike play. For example, upon hearing your proposal, I might not have been at all uncertain about what to do; straight-away, I nonactionally acquired an intention to join you, and I voiced that intention in an overt action of assenting to your proposal. (In that case, what happens in me is like what happens in Helen in one of the beer scenarios discussed earlier.) Or I might have weighed the pros and cons, decided that it would be best to join you, and, on the basis of that cognitive decision, nonactionally acquired an intention to join you. However, there seems also to be a distinctly different possibility. Perhaps, because I already had plans and because your offer was attractive, I was uncertain about what to do. Perhaps, upon reflection, I judged that I could revise my plans without much inconvenience but was still uncertain about what to do, since my prior plans were attractive as well. And perhaps I performed a mental action of assenting to your proposal and then expressed that inner assent to you. In performing that

mental action, if that is what occurred, I *decided* to join you: my mentally assenting to your proposal was an action of intention formation, an action of settling on joining you to watch Mike play tonight.

My own theoretical proposal about deciding to act is that it is a mental action of executive assent to a first-person plan of action. In deciding to act, one forms an intention to act, and in so doing one brings it about that one has an intention that incorporates the plan to which one assents. In some cases, the plan is very schematic; details need to be filled in. For example, without having given any thought to the details, Ann may decide now to vacation in Hawaii for the first time two summers from now. She can settle the details later, at her leisure. In other cases, the details have been worked out, and one has finally resolved one's uncertainty about whether to execute the detailed plan. (Ann may have formulated a specific plan for proposing marriage to Bob while being uncertain about whether to propose marriage.) There is, of course, a range of degrees of specificity in between.

In the case of a decision *not* to *A*, where not *A*-ing is represented by the agent as a genuine not-doing, the plan associated with the intention formed is simply *not to A*. Such an intention is not, in my view, an intention to act (Mele 2003a, pp. 146–54). Intentions to see to it that one does not *A*, however, are intentions to act. And one's having an intention not to *A* may help to account for one's acquiring an intention to see to it that one does not *A* when, for example, one believes that one is in danger of succumbing to temptation to *A*.

My proposal requires that there be some plan that one assents to in making a practical decision (the limiting case for decisions to act, again, being an action-plan with a single node). That is as it should be. Assuming that we decide for reasons, notice that at least minimal representations of a relevant sort are present in reasons for which we decide, on a Davidsonian view of reasons.¹⁷ If Al decided to mow his lawn this morning because he wanted to repay his neighbor for rudely awakening him and believed that his mowing his lawn this morning would constitute suitable repayment, his belief includes a representation of his mowing his lawn this morning—what he decided to do. In deciding to mow his lawn this morning, Al brings it about that he has an intention that incorporates that representation. However, that representation need not exhaust the representational content of the intention that Al forms. A first-person lawn-mowing schema may be stored in Al's memory, and that schema may be part of the content of the intention he forms in deciding to mow his lawn this morning. Some practical decisions to *A* may integrate (some of) the representational content of the reason or reasons for which one decides to *A* with other representational content.¹⁸

Practical deciding is a direct way of resolving practical uncertainty. Of course, someone's deciding to *A* may generate new practical uncertainties. Now that Ann has decided to spend some time in Hawaii two years from now, she is uncertain about how she will get there, where she will stay, and so on. These matters may be settled eventually by future practical decisions. Similarly, in some circumstances, agents who decide not to *A* subsequently decide what to do instead.

Again, if there are basic actions, practical decidings, on the present construal, are among them. Practical decidings, then, are relatively simple structurally; their complexity is in the content of the intention formed, not in the structure of the action of decision-making itself. Deciding to *A* is more like raising one's arm than like signaling that one wants to ask a question by raising one's arm. The proposal that practical deciding is a momentary mental action of executive assent to a first-person plan of action (or, in some cases, a momentary action of settling on not *A*-ing) coheres with this relative simplicity.

Early in this chapter, I asked how practical deciding differs from nonactionally acquiring an intention. If what I have said about practical deciding is correct, there is a correct answer. Practical deciding is a mental action of the sort I have been discussing, and nonactionally acquiring an intention plainly is not an action of any kind.

We make practical decisions only when we are uncertain or unsettled about what to do.¹⁹ It is also true that, normally, at least, when we reason about what it would be best, or "good enough," or permissible to do, we do so in situations involving practical uncertainty. Again, our practical decisions may accord with cognitive decisions that we have reached, or conflict with them, as in some cases of akratic action. If an agent who judges it best to *A*, *decides* to *A*, or *decides* not to *A*, his practical decision resolves practical uncertainty that his cognitive conclusion left unresolved. I have suggested that some cognitive decisions resolve practical uncertainties in a less direct way: if an agent is uncertain about whether to *A*, his acquiring the belief that it would be best to *A* may straightaway result in his acquiring an intention to *A*, and in acquiring that intention he becomes settled on *A*-ing (thereby removing the relevant practical uncertainty).²⁰ But I have also suggested that some practical uncertainties—including some uncertainties that any cognitive decision one reaches fails to resolve—are resolved instead by practical decisions. Practical deciding is an important mode of uncertainty resolution, and it is a straightforwardly actional mode—the mode being momentary mental action of intention

formation, as opposed, for example, to the extended process featured in view 2. That we sometimes decide, in this actional sense, to do things is a pronouncement of ordinary experience, as I explained in section 2.1; and I take section 2.2 to have resolved the strongest worry about the truth of that pronouncement.²¹

This chapter has had a pair of aims: to clarify what practical deciding is, and to defend the idea that there are genuine instances of practical deciding, on the actional account of the phenomenon that I have advanced. In connection with the second aim, I have been willing to be very receptive to claims that what we find in a variety of cases is nonactional intention acquisition rather than actions of practical decision-making. (Again, I made and defended such claims myself in Mele 1992a, chap. 12; but my concern now is the dialectical situation in this chapter.) My contention is that, even on the assumption that nonactional intention acquisition is very common, there is good reason to believe that not all intentions are nonactionally acquired and that there are actions of practical decision-making. Should it turn out that I have been *overly* receptive to nonactional intention acquisition—should it be true, for example, that whenever cognitive decisions that it is best to *A* issue smoothly in intentions to *A*, they in fact do so by prompting a mental action of deciding to *A*, or that seeing something one likes in fact leads one to try to acquire it only if it leads one to perform the mental action of deciding to acquire it—that is not a problem for the view advanced here.²² The upshot would be that practical deciding, on the view of it that I have advanced, is more pervasive than I have been willing to claim it is. That obviously is consistent with my having achieved the aims I identified.

The power to make practical decisions lies at the heart of much of the literature on free will—incompatibilist and compatibilist literature alike. Agents who act intentionally but lack this power are conceivable (Frankfurt 1988, p. 176). Perhaps cats and dogs are like that. They may have beliefs and desires, nonactionally acquire relatively short-term intentions on the basis of their beliefs and desires, and execute those intentions in intentional actions. This is consistent with their never *forming* an intention. We would like to believe that we do form intentions, that we make up our minds—or decide—to do certain things and not to do others. If I am right, it is reasonable to believe this. Setting aside practical decisions for not-doings, it is reasonable, as well, to understand practical decisions as mental actions of executive assent to first-person plans of action. Again, in the case of simple not-*A*-ings, the plan assented to in deciding not to *A* is *not to A*.²³

Notes

1. See Frankfurt 1988, pp. 174–76; Kane 1996, p. 24; Kaufman 1966, p. 34; McCann 1986a, pp. 254–55; Mele 1992a, p. 156; Pink 1996 p. 3; and Searle 2001, p. 94.
2. Brian O’Shaughnessy defends a view of this kind (1980, vol. 2, pp. 300–301). Also see Williams 1993, p. 36.
3. On akratic failures to intend, see Audi 1979, p. 191; Davidson 1980, chap. 2, 1985, pp. 205–6; Mele 1992a, pp. 228–34, and 2012a, pp. 25–28; and Rorty 1980.
4. Incidentally, as I see it, to say that one was up all night deciding *whether* to major in English—or deciding what major to declare—also is to speak loosely. What one means, I think, is that one was up all night deliberating about whether to major in English—or deliberating about what major to declare—and finally settled the matter.
5. Jeff Miller and Wolf Schwarz report that although I reserve “the term ‘decision’ for the final resolution at the end of” a process, they “prefer to use ‘decision’ as shorthand for the entire process rather than reserving it for the final termination” (2014, p. 18). (They refer to this process as “the decision-making process.”) What Miller and Schwarz represent as shorthand may be regarded by some readers as ordinary usage. Some such readers may distinguish between momentary actions of intention formation and relevant processes that lead up to and include such actions, contend that things of both kinds exist, and refer to things of both kinds as “decisions.” Others may contend that certain processes are decisions and deny that there are momentary actions of intention formation. Obviously, my primary concern is with the phenomena—not ordinary usage; and I certainly do not deny (as Miller and Schwarz observe, 2014, p. 18) that decisions are associated with processes that issue in them.
6. A predictable positive effect of recent metaphilosophical attention to intuitions and their place in philosophy is increased caution in first-order philosophy about how one uses the term “intuition,” a term that is used in a variety of different ways by philosophers (see Cappelen 2012). In this book, I use the term “intuition” specifically in connection with reactions to scenarios. In this context, what I have in mind are beliefs and inclinations to believe that are relatively pre-theoretical. These beliefs and inclinations are not arrived at by consulting one’s favorite relevant philosophical position and applying it to the case at hand, and they sometimes prove useful in testing philosophical analyses or theories by testing their implications about cases. I definitely do not regard intuitions as the final word. We may question, test, and reject our own intuitions about cases. I also have no wish to tell others how they should use the word “intuition.”
7. Being uncertain about what to do should not be confused with not being certain about what to do. Rocks are neither certain nor uncertain about anything.
8. McCann’s positive view is that intending to decide to *A* is a constituent of deciding to *A*. For criticism, see Mele 1997, pp. 242–43.
9. On this, McCann and I agree. Intending to *A*, as I understand intentions, encompasses being settled on *A*-ing (Mele 1992a, chap. 9), and this settledness is a

psychological commitment to *A*-ing. Intentions not to *A* encompass settledness on not *A*-ing. (Like Bratman [1987, p. 5] and others, I take this kind of commitment to be revocable. Acquiring an intention does not irrevocably stick one with it. We can change our minds.)

10. In Gregory Kavka's well-known toxin puzzle (1983), there is a payoff specifically for *intending* to drink a certain toxin—a reward for intending to drink it that is not contingent on one's drinking it. The payoff might just as well have been for *deciding* to drink the toxin.
11. What I say here is at odds with what Michael Bratman (1984) has dubbed "the simple view"—the thesis that *SA*-ed intentionally only if *S* intended to *A* (see Adams 1986; McCann 1986b, 1991). For detailed criticism of the simple view, see Bratman 1987, chap. 8; also see Mele 1992a, chap. 8.
12. In Mele 2000, on which this chapter is based, this section was followed by a section examining the worry that practical decisions are inexplicable and their inexplicability counts against their existence. The omitted section seemed outdated.
13. I limit my discussion of intentions here to occurrent intentions (see Mele 2007b, from which the next several paragraphs derive). To avoid confusion, two different tendencies in the literature on the connection between plans and intentions should be mentioned. The contrast is nicely illustrated by Myles Brand's claim that "the cognitive component of prospective intention is a plan" (1984, p. 153) and Michael Bratman's assertion that "We form future-directed intentions as *parts* of larger plans . . ." (1987, p. 8; my emphasis). The difference is largely a matter of convention. Although I believe that intentions are sometimes formed with a view to the execution of larger plans (that is, plans larger than the ones embedded in the intentions at issue), I have not adopted Bratman's convention of using the word "plan" for items that themselves involve "an appropriate sort of commitment to action" (1987, p. 29).
14. The guidance depends on the agent's monitoring progress toward his goal. The information (or misinformation) that Jan has entered her password, for example, figures in the etiology of her continued execution of her plan. On guidance, see Mele 2003a, 55–62.
15. In the case of an intention for a not-doing (e.g., an intention not to vote tomorrow), the agent may instead be settled on not violating the simple plan embedded in it—the plan not to vote. On not-doings and attitudes toward them, see Mele 2003a, 146–54.
16. A critic may claim that in all cases of this kind the agent is settled on a course of action without realizing it and that he is deliberating only because he does not realize what he is settled on doing. For argumentation to the contrary, see Mele 1992a, chap. 9.
17. See McCann 1998, chap. 8. On some other views of reasons, representations of the sort I will identify are not present in the reasons for which we decide.

18. On the representational content of intentions, see Mele 1992a, chap. 11. Obviously, I am not suggesting that all reasons for which we decide are represented in the contents of our decisions. See Mele 2003a, pp. 42–45.
19. There are very special cases in which an agent who is offered a reward for deciding to *A* is convinced that he will *A* whether or not he intends (or decides) to *A* (Mele 1992b). Perhaps such an agent may decide to *A*. An agent offered Kavka's (1983) toxin deal may be convinced that someone will cause him to drink the toxin unintentionally if he does not intentionally drink it. When offered a prize for deciding to drink the toxin, he may be uncertain about whether to make this decision. He may be uncertain about that even though—being convinced that if he does not drink the toxin intentionally, he will drink it unintentionally—he is not uncertain about whether he will drink the toxin.
20. In Mele 1995a, pp. 25–30, I attempt to explain why “best judgments” sometimes result in corresponding intentions and sometimes fail to do so.
21. Is there neuroscientific evidence of the existence of practical decisions, as I conceive of them? A study by Marjan Jahanshahi and colleagues compares brain activity in subjects who are following the instruction to raise their right index finger whenever they wish with brain activity that occurs when these subjects are instead following the instruction to raise their right index finger whenever they hear a tone (Jahanshahi et al. 1995). Subjects in the former condition might *decide* when to raise their finger, whereas subjects in the latter condition are simply raising it in response to a tone—intentionally, of course. Jahanshahi et al. found greater activation of the dorsolateral prefrontal cortex in the former condition, and they infer that it “was associated with the additional requirement of decision-making about the timing of the movement on each trial, or ‘when to do’ it” (1995, p. 930; my emphasis). We may have here physical evidence of a difference between proximally deciding to *A* and otherwise acquiring a proximal intention to *A*, if detection of the tone prompts a proximal intention to raise the finger (but for a relevant caveat, see Jahanshahi et al. 1995, p. 930). For evidence that activity in the presupplementary motor area is associated specifically with deciding, see Lau et al. 2004.
22. A theorist who holds that cognitive and perceptual events of these kinds can issue in intentions or actions only by way of practical decisions regards such cognitive and perceptual events as incapable of resolving practical uncertainty.
23. It is sometimes claimed that scientific findings warrant the claim that we never *consciously* make decisions. For a detailed rebuttal of this claim about scientific findings, see Mele 2009. For comments on a draft of Mele 2000, on which this chapter is based, I am grateful to Bruce Aune, Randy Clarke, John Heil, and Hugh McCann.

Actions, Explanations, and Causes

WHAT CONDITIONS MUST an adequate explanation of any intentional human action satisfy? According to *causalists* about action explanation, one necessary condition is that the explanation cite a cause of the action. *Anticausalists* about action explanation disagree. In this chapter, I explore various options causalists have for a required causal condition and I critically examine leading anticausalist proposals.

3.1. Causalism and Reasons

Here is one candidate for a causalist proposal:

Dr. Necessarily, if E is an adequate explanation of an intentional action A performed by an individual agent S , E cites a reason that was a cause of A .

(Once again, I leave it to readers to individuate actions as they deem best and to read “ A ” as a variable either for actions themselves or for actions under A -descriptions, depending on their preferred mode of act-individuation.) According to one way of thinking about reasons for action, advocated by Donald Davidson (1980, chap. 1) and others, they are composed of beliefs and desires. For example, Don’s reason for flipping the switch might have been composed of a desire to have sufficient light for reading and a belief that flipping the switch would produce such light. Some philosophers have rejected this conception of reasons for action (see, for example, Dancy 2000, Scanlon 1998). Imagine a philosopher who holds that the reasons for which agents

act are limited to true propositions and that no proposition can be a cause of anything. Such a philosopher obviously will reject *DI*. Later in this section, I consider a variant of *DI*—one featuring belief, desire, and intention—that may, in principle, win the approval of the imaginary philosopher. Some background is in order first.

As I observed elsewhere (Mele 2003a, p. 69), philosophical work on what its authors call reasons for action tends to be guided by concerns with two distinct but related topics: the *explanation* of intentional actions; and the *evaluation* of intentional actions or their agents. In work dominated by the explanatory concern, reasons for action tend to be understood as states of mind (for example, as certain kinds of combinations of beliefs and desires à la Davidson). In some work dominated by the evaluative concern, typical reasons for action are understood as states of—or facts or true propositions about—the agent-external world.

Jonathan Dancy writes that

intuitively it seems to be not so much propositions as states of affairs that are our good reasons. It is her being ill that gives me reason to send for the doctor, and this is a state of affairs, something that is part of the world, not a proposition. Those who announce that all good reasons are propositions (e.g. Scanlon 1998: 57) seem thereby to lose contact with the realities that call for action from us. (2000, pp. 114–15)

According to what I elsewhere dubbed an *objective favorers* view of reasons for action, typical reasons for action are agent-external states of affairs (Mele 2007c). On an *exclusivist* version of this view, the only reasons for action are items that objectively favor the actions for which they are reasons. (Some states of an agent may objectively favor an action: for example, Joe's belief that demons dance in his kitchen may objectively favor his seeking psychiatric help [Mele 2007c, p. 91; see Dancy 2000, p. 124].)

Suppose that Dave, a causalist with Davidsonian leanings, were to agree, at least for the sake of argument, to adopt an exclusivist objective favorers view of reasons for action. What might Dave say about *DI*? Dave might consider holding on to *DI* and trying to develop a position according to which reasons (objective favorers) are always among the causes of beliefs that in turn are among the causes of actions. For example, if Ann's being ill was Bob's reason for phoning a doctor, Dave might say that her being ill was among the causes of Bob's acquiring a belief that she is ill and that (his acquiring) that belief was among the causes of his phoning the doctor.

However, it should occur to Dave that some intentional actions are not objectively favored by anything at all. For example, Nick, who believed Angela to be at home, drove to her house to help her assemble a bookcase. Unfortunately, Angela was unexpectedly called away before Nick arrived. We might say that something *subjectively* favored Nick's driving to Angela's house—perhaps the combination of his desire or intention to help Angela and his belief that he would put himself in a position to do that by driving to her house. But, as it turned out, nothing *objectively* favored it. Even so, Nick's driving to Angela's house—an intentional action—certainly seems to be explicable. It is explicable even though there is no objective favorer—and hence no reason, on the view at issue—to be a cause of it. Dave realizes that given his agreement to speak as an exclusivist objective favorers theorist, he should abandon *DI*.

Is there a reasonable candidate for a replacement? Dave considers the following:

D2. Necessarily, if *E* is an adequate explanation of an intentional action *A* performed by an individual agent *S*, *E* cites a belief that was a cause of *A*.

Dave notices that the stories about Bob and Nick have something interesting in common. Bob believes that Ann is ill, and Nick believes that Angela is at home and will still be there when he arrives. Partly because Ann is ill, something objectively favors Bob's phoning a doctor; but because Angela is not at home, nothing (other things being equal) objectively favors Nick's driving to her house. Even so, both agents act as they do partly because they believe what they do. In Bob's case, an objective favorer might be among the causes of a pertinent belief that is a less remote cause of his phoning the doctor; and in Nick's case, there is no objective favorer to play a causal role of this kind. But Dave finds the thought that, in both cases, a relevant belief seems to be doing significant work reassuring.

Dave has learned to be cautious. He has heard of wholly intrinsically motivated actions (see Mele 1992a, pp. 104–12)—actions done solely for their own sakes. And he believes, for example, that it is possible for someone to whistle a tune—intentionally—for no further purpose at all. Dave does not understand the notion of reasons as objective favorers well enough to be confident whether when someone whistles a happy tune simply because he feels like it, as one might say, his action is likely to be objectively favored by something. Nor does he see any reason to insist that explaining such an action requires

citing a belief that was a cause of it (see Mele 1992a, chap. 6). So he opts for something more cautious than *D2*:

D3. Necessarily, if *E* is an adequate explanation of an intentional action *A* performed by an individual agent *S*, *E* cites a belief, desire, or intention that was a cause of *A*.

The main point to be emphasized now is that even if Davidson is wrong about what reasons for action are, that is compatible with the truth of causalism about action explanation. Indeed, even if all reasons for action are true propositions and true propositions cannot be among the causes of anything, it cannot be inferred from this that causalism about action explanation is false. The suppositions about reasons and propositions just identified are compatible with the supposition that *D3* is true, and *D3*'s truth is sufficient for the truth of causalism about action explanation.¹

3.2. Causalism and States of Mind

Some philosophers posit token mental states such as intentions, desires, and beliefs and attribute causal roles to these states or to their neural realizers in the production of actions (Brand 1984; Davidson 1980, chap. 1; Mele 1992a, 2003a). Other philosophers are wary of postulating such token states of mind (Child 1994; Hornsby 1993). They would not accept *D3*. Suppose that although people believe things, desire things, and intend to do things, there are no token states of mind—no beliefs, no desires, and no intentions, for example. Would that falsify causalism about action explanation?

Perhaps not. Fact causation is an option (Mellor 1995). And because it is, so is the following claim:

D4. Necessarily, if *E* is an adequate explanation of an intentional action *A* performed by an individual agent *S*, *E* cites a fact about something the agent believed, desired, or intended, which fact was a cause of *A*.

A cautious causalist about action explanation can opt for the following disjunctive claim:

D5. Necessarily, if *E* is an adequate explanation of an intentional action *A* performed by an individual agent *S*, then *E* cites (1) a reason that was

a cause of *A* or (2) a belief, desire, or intention that was a cause of *A* or (3) a neural realizer of a belief, desire, or intention, which neural realizer was a cause of *A* or (4) a fact about something the agent believed, desired, or intended, which fact was a cause of *A*.

*D*₅'s truth is sufficient for the truth of causalism about action explanation.

3.3. Causalism and Causation

Teleological explanations of human actions are explanations in terms of aims, goals, or purposes of human agents. Some proponents of the view that human actions are explained teleologically regard all causal accounts of action explanation as *rivals* (Sehon 1994, 1997, 2005; Taylor 1966; Wilson 1989, 1997). Scott Sehon asserts: "Teleological explanations simply do not purport to be identifying the cause of a behavior" (2005, p. 218). But, as David Lewis observes, speaking in terms of "the cause of something" can easily generate confusion (1986, p. 215). Lewis adds: "If someone says that the bald tire was the cause of the crash, another says that the driver's drunkenness was the cause, and still another says that the cause was the bad upbringing which made him so reckless, I do not think any of them disagree with me when I say that the causal history includes all three." In any case, causalists like me do not purport to be identifying *the cause* of an action when we offer causal explanations of actions in terms of agents' aims, goals, or purposes. The basic idea—oversimplifying a bit—is that a putative teleological explanation of an action in terms of a goal, aim, or purpose *G* does not explain the action unless the agent's wanting or intending to (try to) achieve *G* has a relevant *effect* on what he does. (My expression "the agent's wanting or intending to (try to) achieve *G*" is meant to leave open both state readings and fact readings.) Obviously, the notion of *having an effect* is a causal notion; and the assertion, for example, that an agent's intending to achieve *G* had an effect on what he did places his intending to do that in the causal history of what he did.

Lewis defends the thesis that "to explain an event is to provide some information about its causal history" (1986, p. 217). Here is a more modest thesis that is entailed by Lewis's thesis: (*C*₁) Nothing is an explanation of an event unless it provides some information about the event's causal history. (The claim that *E* had no causes—or no causal history—is not counted as a claim that provides information about *E*'s causal history.) If *C*₁ is true, then if all intentional actions are events, the following thesis also is true: (*C*₂) nothing

is an explanation of an intentional action unless it provides some information about that action's causal history. Now, a theorist who holds that there are both adequate and inadequate explanations may, in principle, contend both that C_2 is false and that C_2 would be true if "explanation" were modified by "adequate." To keep things simpler than they would otherwise be, I will often use "explanation" as shorthand for "adequate explanation." C_2 is a statement of causalism about action explanation, and it is a consequence of Lewis's thesis about event explanation in general, on the assumption that all actions are events.

Anticausalists about action explanation who take all actions to be events are committed to denying Lewis's thesis about event explanation. Perhaps a source of part of this disagreement is a difference in how causation is conceived. In the literature defending or attacking causalism about action explanation, not much is said about what causation is. This is understandable: both sides may be thinking that, whatever the best account of causation is, their view about how actions are to be explained is the correct one to take, and they may want to avoid hitching their wagon to a specific theory of causation. Possibly, differences in how some causalists and anticausalists about action explanation understand causation help account for their different positions on whether causalism is true or false. Precisely because the body of literature at issue does not have much to say about the nature of causation, the possibility just identified is difficult to explore.

3.4. A Challenge for Anticausalists

Donald Davidson writes that "a person can have a reason for an action, and perform the action, and yet this reason not be the reason why he did it" (1980, p. 9). Two pages later, he asserts that "failing a satisfactory alternative, the best argument for a [causal] scheme like Aristotle's is that it alone promises to give an account of the 'mysterious connection' between reasons and actions" (p. 11). These two remarks lie at the heart of an important challenge to anticausalists.

I set the stage for a version of Davidson's challenge with a story. Two different things, T_1 and T_2 , independently dispose Al to mow his lawn this morning. T_1 has to do with schedule-related convenience and T_2 with vengeance. Al wants to mow his lawn this week and he believes that this morning is a convenient time, given his schedule for the week. But he also wants to repay his neighbor for the rude awakening he suffered recently when she turned on her mower at the crack of dawn, and he believes that mowing his lawn this

morning would constitute suitable repayment. As it happens, Al's purpose in mowing his lawn this morning accords with one or the other of T_1 and T_2 but not both.

Suppose that Al's wanting to mow his lawn this week and his believing what he does about convenience have no effect at all on what Al does with his lawnmower this morning. Might it be the case, even so, that he mows his lawn for a purpose or reason that accords with T_1 rather than for a purpose or reason that accords with T_2 ? A philosopher who answers in the affirmative faces the challenge of producing an informative conceptually sufficient condition for "Al mows his lawn this morning for a purpose or reason that accords with T_1 " that is consistent with this answer and the details of the story.

I told Al's story in a way that is noncommittal about the existence of token states of mind, and my expressions "Al's wanting to mow his lawn this week" and "his believing what he does about convenience" are meant to leave open both *state* and *fact* readings. Believers in token states of mind can ask themselves whether Al mows his lawn for a purpose or reason that accords with T_1 rather than for a purpose or reason that accords with T_2 even though his belief about convenience and his desire to mow his lawn this week (and their neural realizers) are not among the causes of his mowing. And those who are skeptical about the existence of token states of mind can ask themselves a parallel question in terms of the fact that Al believed what he did about convenience and the fact that he wanted to mow his lawn this week.

One way—a very effective way, I believe—to give readers a firm sense of how challenging Davidson's challenge (or a variant thereof) is to show that the detailed attempts anticausalists have made to meet it fail. That is the business of the remainder of this section.

Two ways of evading the challenge should be mentioned at this point. It may be claimed that whenever agents who A intentionally have two or more reasons for A -ing, (1) they A for all of them, or, alternatively, (2) there is no fact of the matter about which are the reasons for which they A . Not only are both contentions ad hoc, both also are challenged later in this chapter.

3.4.1. *Wilson 1989 and Sehon 1994*

After summarizing Davidson's challenge (1989, pp. 168–70), George Wilson complains that anticausalists have left us pretty much in the dark about the connection between reasons and actions and he remarks that "without some more positive identification of the putatively noncausal connection in question, the nature of this linking remains mysterious" (p. 171). Wilson offers

a detailed reply to the challenge. Regarding a particular case, he claims that facts of the following kind “about the context of the action, about the agent’s perception of that context, and about the agent’s sentient relations to his own movements . . . make it true that he intended of those movements that they promote his getting back his hat [read: make it true that those movements were *sentiently directed* by him at promoting his getting back his hat]”:²

[1] The man, wondering where his hat is, sees it on the roof, fetches the ladder, and immediately begins his climb. [2] Moreover, the man is aware of performing these movements up the ladder and knows, at least roughly, at each stage what he is about to do next. [3] Also, in performing these movements, he is prepared to adjust or modulate his behavior were it to appear to him that the location of his hat has changed. [4] Again, at each stage of his activity, were the question to arise, the man would judge that he was performing those movements as a means of retrieving his hat. (1989, p. 290)

Here Wilson apparently offers what he takes to be conceptually sufficient conditions for its being true that the man sentiently directed certain movements of his at promoting his getting back his hat. Are these conditions in fact sufficient for this?³

Two points need to be made before this question is answered. The first concerns the second sentence of the passage just quoted. What knowledge is attributed to the man there? Not the knowledge that he is about to perform a sentiently directed movement of a certain kind, if sentient direction is a notion that this passage is supposed to explicate. However, this does not block the supposition that the man knows that he is about to perform a movement of his left hand onto the next rung, for example, in Wilson’s broad sense of “perform a movement” (according to which “a man performs a convulsive and spasmodic movement when he clutches and cannot loose a live electric wire, and someone undergoing an epileptic seizure may perform a series of wild and wholly uncontrollable movements” [p. 49]).

Second, Wilson claims that “to try to [*A*] is (roughly and for the pertinent range of cases) to perform an action that is intended to [*A*]” (p. 270). Brief attention to trying will prove useful. Suppressing Wilson’s qualification for a moment, his claim amounts, for him, to the assertion that to try to *A* is to perform an action that is sentiently directed by the agent at *A*-ing. Regarding the following plausible proposition, there is no need to worry about limiting the range of cases: (*Ti*) if one is doing something that one is sentiently

directing at *A*-ing, then one is trying to *A*, in an utterly familiar, unexacting sense of “trying.” Although people may often reserve attributions of trying for instances in which an agent makes a considerable or special effort, this is a matter of conversational implicature and does not mark a conceptual truth about trying.⁴ A blindfolded, anesthetized man who reports that he has raised his arm as requested, quite properly responds to the information that his arm is strapped to his side by observing that, in any case, he *tried* to raise it—even though, encountering no felt resistance, he made no *special* effort to raise it.⁵ And when, just now, I typed the word “now,” I tried to do that even though I easily typed the word. Now, *T₁* entails (*T₂*) that one who is not trying to *A*, even in the unexacting sense of “trying” identified, is not sentiently directing one’s bodily motions at *A*-ing. *T₂* is very plausible, as is the following, related proposition: (*T₃*) one who is not trying to do anything at all, even in the unexacting sense of “trying,” is not sentiently directing one’s bodily motions at anything.

In the following case, as I will explain, although Norm is not, during a certain time, trying to do anything, even in the unexacting sense of “trying” that I identified, he satisfies all four of the conditions in the quotation at issue during that time. The moral is that these conditions are not sufficient for a person’s sentiently directing “movements” of his at the time. Their insufficiency follows from the details of the case and the platitude that an agent who is not trying (even in the unexacting sense) to do anything is not sentiently directing his bodily motions at anything.

Norm has learned that, on rare occasions, after he embarks on a routine activity (for example, tying his shoes, climbing a ladder), Martians take control of his body and initiate and sustain the next several movements in the chain while making it seem to him that he is acting normally. He is unsure how they do this, but he has excellent reason to believe that they are even more skilled at this than he is at moving his own body, as, in fact, they are. (The Martians have given Norm numerous demonstrations with other people.) The Martians have made a thorough study of Norm’s patterns of peripheral bodily motion when he engages in various routine activities. Their aim was to make it seem to him that he is acting while preventing him from even *trying* to act by selectively shutting down portions of his brain. To move his body, they zap him in the belly with M-rays that control the relevant muscles and joints. When they intervene, they wait for Norm to begin a routine activity, read his mind to make sure that he plans to do what they think he is doing (for example, tie his shoes or climb to the top of a ladder), and then zap him for a while—unless the mind-reading team sees him abandon or modify his

plan. When the mind readers notice something of this sort, the Martians stop interfering and control immediately reverts to Norm.

A while ago, Norm started climbing a ladder to fetch his hat. After he climbed a few rungs, the Martians took over. Although they controlled Norm's next several movements while preventing him from trying to do anything, they would have relinquished control to him if his plan had changed (for example, in light of a belief that the location of his hat had changed).

Return to facts 1 through 4. Fact 1 obtains in this case. What about fact 2? It is no less true that Norm performs his next several movements than that the man who clutches the live electric wire performs convulsive movements. And the *awareness* of performing movements mentioned in fact 2 is no problem. The wire clutcher can be aware of bodily "performances" of his that are caused by the electrical current, and Norm can be aware of bodily "performances" of his that are caused by M-rays. Norm also satisfies a "knowledge" condition of the sort I identified. If Wilson is right in thinking that an ordinary ladder climber knows, in some sense, that he is about to perform a movement of his left hand onto the next rung, Norm can know this too. What he does not know is whether he will perform the movement on his own or in the alternative way. But that gives him no weaker grounds for knowledge than the ordinary agent has, given that the subject matter is the performance of movements in Wilson's broad sense and given what Norm knows about the Martians' expertise. Fact 3 also obtains. Norm is prepared to adjust or modulate his behavior. (And it is possible for him to do so. Although the Martians in fact initiated and controlled Norm's next several movements up the ladder while preventing him from trying to do anything, they would not have done so if his plans had changed.) Fact 4 obtains too. In Wilson's sense of "perform a movement," Norm believes that he is performing his movements "as a means of retrieving his hat." He does not believe that the Martians are controlling his behavior; after all, he realizes that they very rarely do so.

Even though these facts obtain, Norm does not sentiently direct his next several movements up the ladder at getting his hat because he is not sentiently directing these movements at all. Wilson maintains that sentiently directing a bodily movement that one performs entails exercising one's "mechanisms of . . . bodily control" in performing that movement (1989, p. 146). However, Norm did not exercise these mechanisms in his performance of the movements at issue. Indeed, he did not make even a minimal effort to perform these movements; owing to the Martian intervention, he made no effort at all—that is, did not try—to do anything at the time. And it is a platitude that one who

did not try to do anything at all during a time t did not sentiently direct his bodily motions during t .

It might be suggested that although Norm did not directly move his body during the time at issue, he sentiently directed his bodily motions in something like the way his sister Norma sentiently directed motions of her body when she orally guided blindfolded colleagues who were carrying her across an obstacle-filled room as part of a race staged by her law firm to promote teamwork. If Norma succeeded, she may be said to have brought it about that she got across the room, and her bringing this about is an action.⁶ Notice, however, that there is something that she was trying to do at the time. For example, she was trying to guide her teammates. By hypothesis, there is nothing that Norm was trying to do at the relevant time, for the Martians blocked brain activity required for trying. And this is a crucial difference between the two cases. The claim that Norma sentiently directed motions of her body at some goal at the time is consistent with T_3 ; the comparable claim about Norm is not.⁷

Wilson proposed sufficient conditions for its being true that a person's movements were sentiently directed by him at promoting his getting back his hat. Norm satisfies those conditions even though it is false that the "movements" at issue were sentiently directed by him. So those conditions are not in fact sufficient.

Can Wilson's proposal be rescued simply by augmenting it with an anti-intervention condition? No. If the addition of such a condition does contribute to conceptually sufficient conditions for a person's sentiently directing his movements at a goal, it may do so because the excluded kinds of intervention prevent, for example, the obtaining of normal causal connections between mental items or their neural realizers and bodily motions. An anticausalist who augments Wilson's proposal with an anti-intervention condition also needs to produce an argument that the condition does not do its work in this way.

I turn to Scott Sehon's 1994 attempt to answer Davidson's challenge.⁸ Under the heading "Defusing the Davidsonian Challenge," he argues that a teleologist can appeal to counterfactuals to "distinguish between reasons an agent acted on and reasons the agent had but did not act on" (p. 67). Sehon invites us to imagine that Heidi "lifts a heavy book up to the top of a bookshelf" while having the following pair of desires: "a desire to put the book where it belongs and a desire to strengthen her biceps." He assumes that only one of these desires "provides the reason why Heidi lifted the book," and he asks which one does so. Sehon reports that this question, "as viewed from

the teleological theory, looks roughly like this: toward which outcome did Heidi direct her behavior, that the book was put away or that her biceps were strengthened?" The correct answer, he contends, depends on what counterfactuals of a certain kind are true of Heidi at the time. If the book had belonged on the bottom shelf, would she have put it there, or would she have placed it on the top shelf? If something more suitable for the purposes of exercising her biceps had been present, would she have lifted it, or would she have lifted the book? And so on.

Suppose that the counterfactual test indicates that Heidi's goal was that the book be returned to its proper place. Even so, one will not know in virtue of what it is true that Heidi directed her behavior toward that goal until one knows in virtue of what it is true that Heidi directed her behavior. One can apply Schon's counterfactual test to a case in which Martians who wished to deceive Heidi into believing that she was acting manipulated her muscles in order to make her bodily motions fit the intention she had at the time to return the book to its proper place while preventing her from even trying to return it, and one would get the result that Heidi directed her behavior toward the book's being returned there. But, of course, that result would be false, since Heidi was not directing her behavior—that is, acting—at all in this case. Rather, the Martians were controlling the motions of her body.

True counterfactuals are true in virtue of something or other. Their truth is grounded in something factual. If Heidi was executing—that is, acting on—an intention to return the book to its proper place, then, other things being equal, one should expect such counterfactuals as the following to be true: if Heidi had believed that the book belonged on the bottom shelf, she would have placed it there; if Heidi had believed that the book's proper place was the middle shelf, she would have put it there. But if these counterfactuals are true for the reasons one expects them to be, their truth is grounded in part in Heidi's *acting* with the intention of putting the book where it belongs; their truth does not explain what it is for Heidi to be acting with this intention.

One moral of the objections I have raised to Wilson's and Schon's attempted answers to the Davidsonian challenge is clear. Unless an item of one of the kinds featured in *D5* (for example, an intention or its physical realizer) plays a causal role in the production of a person's bodily motions, and not simply the causal role of providing information about goals to mischievous Martians, there is the threat (as in my Martian chronicles) that the person is not acting at all, much less acting in pursuit of the goal(s) that the desire or intention specifies. Partly because teleologists have not offered an acceptable account of what it is to *act*, or to "direct" one's bodily motions,

they have not offered an acceptable account of what it is to act for the sake of a particular goal.

3.4.2. *Sehon 2005 and D'Oro 2007*

As I mentioned in section 3.3, some proponents of the view that human actions are explained teleologically regard all causal accounts of action explanation as rivals. In Mele 2003a (p. 38), I dubbed this position “anticausalist teleologism” (*AT*, for short). In a book defending *AT*, Sehon replies to my objection to Wilson’s proposal (2005, pp. 167–71). He contends that it “is not obvious that Norm’s behavior fails to be an action” (2005, p. 168). In support of this claim, Sehon sketches a version of my story in which Norm is about to shoot someone when “the Martians take over his body and make it carry out the dirty deed.” Sehon asks, “Does this completely absolve Norm of responsibility for shooting his professor?” He reports that it is not obvious to him how to answer this question. “Accordingly,” he writes, “even in the routine case of going up the ladder, I take it not to be obvious that Norm failed to act.”

Things are not nearly as murky as Sehon thinks. I start with an obvious point. Norm is responsible for shooting his professor only if Norm shoots his professor. This is an instance of the truism that a person *P* is responsible for doing *A* only if *P* did *A*. But notice that *P* may have some responsibility for *the shooting of X* even if, because *P* did not shoot *X*, *P* has no responsibility for *shooting X*. For example, if *P* hires—or forces—someone to shoot *X*, then, other things being equal, *P* has some responsibility for the shooting of *X*. In Sehon’s story, the Martians “make it seem to Norm as if he is acting, and if Norm had changed his mind and decided to put the gun down [they] would have immediately relinquished control” to Norm (2005, p. 168). (Here Sehon is following my lead, of course.) So readers may be strongly inclined to see Norm as responsible for more than, in Sehon’s words, merely “having a plan to commit murder.” In light of the distinction I have drawn between *P*’s having some responsibility for shooting *X* and *P*’s having some responsibility for the shooting of *X*, those who, like me, are convinced that Norm did not shoot the professor can maintain that, even so, he has some responsibility for the shooting of the professor.

Imagine that Sehon’s story remains basically the same, except that a Martian makes himself invisible and then, with his slim but powerful tentacle, pulls Norm’s paralyzed finger down on the trigger. (The Martian also makes it seem to Norm as though Norm is pulling the trigger, and “if Norm had changed his mind and decided to put the gun down, the [Martian] would

have immediately relinquished control” to Norm.) Obviously, Norm did not pull the trigger. So Norm did not shoot the professor. That entails that Norm is not responsible for shooting the professor. Even so, he may have some responsibility for the shooting of the professor. After all, the Martian would not have pulled the trigger with Norm’s finger if Norm had not been bent on shooting the professor, and the professor’s life would have been spared if Norm had changed his mind. Schon’s uncertainty about whether Norm shot the professor seems to derive from his failure to distinguish having some responsibility for shooting the professor from having some responsibility for the shooting of the professor.

Some readers may wonder why my story about Norm in Mele 2003a portrays the Martians as using M-rays rather than moving Norm with their tentacles. Recall that Wilson’s conditions include the man’s “performing movements”; and although he is happy to say that a wire clutcher whose body is being jerked about by the wire’s electrical discharge is performing movements in his thin sense, I doubt that he would count a man whose limb motions are caused by Martians pulling and pushing on his limbs as performing movements with those limbs. Naturally, I thought the M-rays were just fine for my purposes, but instead I could have portrayed the Martians as moving Norm’s paralyzed body with just the right sorts of electrical jolts to muscles and joints. Call this *E-manipulation*.

Schon wonders why my Martians interfere with Norm. “What’s in it for the Martians?” he asks (2005, p. 168). In Mele 2003a, I neglected to mention that the Martians had read page 290 of Wilson’s book and wanted to provide a living counterexample to his proposal. In any case, Schon’s reflection on his question leads him to the following claim:

Mele stipulates that the Martians are going to make Norm’s body do exactly what Norm planned to do anyway. If this were an ironclad promise from the Martians, or better yet, something that followed necessarily from their good nature, then . . . I have little problem saying that Norm is still acting, despite the fact that the causal chain involved is an unusual one. If he commits a murder under these circumstances, we will definitely not let him off. (2005, p. 169)

Yes, if Norm commits a murder, he should be blamed for that. But if he is not acting, he commits no murder. Is Norm acting in Schon’s scenario? Presumably, for the purposes of his thought experiment, Schon means to retain as much as he can from my story about Norm while turning the Martians into beings

whose good nature entails that they always make “Norm’s body do exactly what Norm” plans to do. Evidently, Schon is not impressed by the following details of my story: the Martians prevent Norm “from even *trying* to act by selectively shutting down portions of his brain,” and they move his body by zapping “him in the belly with M-rays that control the relevant muscles and joints” (Mele 2003a, p. 49). I am not sure why. Possibly, he rejects T_3 (the thesis that one who is not trying to do anything at all, even in the unexact sense of “trying,” is not sentiently directing one’s bodily motions at anything). And possibly, he accepts T_3 and believes that Norm counts as trying to do things when his Martians replace mine.

Consider a scenario in which, instead of using M-rays, Schon’s good Martians paralyze Norm’s body and then move it by E-manipulation while making it seem to Norm that he is acting normally. When, for example, Norm intends to climb a ladder to get his hat, the Martians paralyze him and E-manipulate his body up the ladder. (They do all this while making it seem to Norm that he is acting normally, and if Norm were to change his mind his paralysis would immediately cease and control would revert to him.) Obviously, in this scenario, Norm is not climbing the ladder. Yet, unless Schon can identify a crucial difference between the use of M-rays and this alternative mode of Martian body manipulation, he is committed to having “little problem saying” that Norm is climbing it.

Schon is willing to grant that when my Martians are at work rather than his, Norm is not acting (2005, p. 168). He contends that, in my story, “since Norm fails . . . to satisfy” the following condition, “his behavior does not count as goal directed” on his “account of the epistemology of teleology” (p. 169): (R_1) “Agents act in ways that are appropriate for achieving their goals, given the agent’s circumstances, epistemic situation, and intentional states” (p. 155). If I am right, Norm is not acting at all, in which case invoking R_1 is overkill. And if Norm is not acting, as Schon is willing to grant, then Wilson’s proposal about sufficient conditions for its being true that a person’s movements were sentiently directed by him at promoting his getting back his hat is false, which is what I set out to show with the Martian example in Mele 2003a.

Some readers may feel that they have lost the plot. The following observation will help. One thing that Schon would like to show is that a proponent of AT can “accommodate our intuition that Norm is not acting” in my story (2005, p. 170). He argues that “Norm’s motion is not that of an agent, because in a range of nearby counterfactual situations his behavior is not appropriate to his goals. Specifically, in all those situations in which the Martians simply

change their mind about what they want to have Norm's body do, Norm's body will do something quite different."

Sehon's explanation of why Norm is not acting is seriously problematic. Imagine a case in which the Martians consider interfering with Norm but decide against doing that. Norm walks to the kitchen for a beer without any interference from the Martians. There are indefinitely many variants of this case in which the Martians change their minds about not interfering and make Norm's body do something else entirely. So "in a range of nearby counterfactual situations his behavior is not appropriate to his goals" (Sehon 2005, p. 170). But this certainly does not warrant the judgment that Norm is not acting in the actual scenario. Obviously, he *is* acting in that scenario: he is walking to the kitchen for a beer. If Sehon is thinking that his counterfactual test for whether an agent is acting is to be applied in scenarios in which the Martians interfere with Norm but not in scenarios in which they do not interfere with him, he does not say why this should be so.

The problems with Sehon's explanation of why it is that Norm is not acting in my case do not end here. He considers a woman, Sally, who "has an odd neurological disorder" (2005, p. 170). When she tries to move her finger in way *W*, her finger often becomes paralyzed and "her body goes through any number of other random motions." In a particular case, Sally successfully tries to move her finger in way *W* when pulling a trigger and murdering a professor. Sehon contends that because Sally's "behavior is generally very sensitive to her goals"—after all, it is "subject to these flukes only when it involves a finger pulling"—she, "unlike Norm, satisfies the condition imposed by (*R*_{*I*}) well enough to make her an agent at the time in question" (p. 171).

This will not do. Imagine a variant of Norm's story in which a rogue Martian interferes with Norm only on one occasion. (The Martian is imprisoned for life by the Martian authorities immediately afterward and no one else ever interferes with Norm.) He moves Norm's paralyzed body up the ladder by E-manipulation while making it seem to Norm that he is acting normally. Sally's behavior is "generally very sensitive to her goals," and I stipulate that Norm's behavior is generally even more sensitive to his goals. Even so, he is not acting as his body moves up the ladder. That Sally is acting whereas Norm is not is not explained by a difference in the general sensitivity of their behavior to their goals.⁹

Sehon concludes his discussion of my objection to Wilson's proposal with the following report:

One could alter the example by making Sally's neurological disorder much more general, such that she rarely does what she intends; but

with that revision, my own intuitions about the case grow flimsy. I'm not sure what to say about her agency in such a case, and I'm not too troubled by the conclusion that she is not exhibiting genuine goal-directed behavior at any particular moment. (2005, p. 171)

Imagine that you read in a medical journal a study of a sixty-year-old patient, Pat, who for the past ten years has suffered from a horrible illness that has the consequence that in about 99% of the cases in which he tries to do something, his "body goes through any number of other random motions" (Sehon 2005, p. 170) and his attempt fails. About 1% of the time, he succeeds in doing what he tries to do. Sometimes, for example, he makes a successful attempt to signal his nurse by pressing a button. Apparently, Sehon would not be "too troubled by the conclusion" that Pat *never* exhibits "genuine goal-directed behavior." But if unwarranted conclusions are troublesome, Sehon should be troubled by this. Pat's story clearly is conceptually possible. In it, some of his behavior is genuinely goal directed.

Two conclusions may now be drawn. First, Sehon has not shown that my objection to Wilson's proposal is unsuccessful. In fact, insofar as he concedes that Norm is not acting in my story, he apparently concedes that the objection is successful. (He nowhere claims that Norm does not satisfy Wilson's proposed conditions.) It is perhaps worth mentioning in this connection that Sehon might have misled not only his readers but also himself by treating my objection to Wilson's proposal as though it were an argument for the claim that no version of *AT* can "accommodate our intuition that Norm is not acting" in my story (Sehon 2005, p. 170). Other anticausalists about action explanation—for example, Ginet and Wallace—have proposed other sufficient conditions for a human being's performing an action or acting in pursuit of a particular goal, and my objections to Ginet's and Wallace's proposals in Mele 2003a (and later in the present chapter) are very different from my objection to Wilson's proposal. Naturally, the objections I offered were designed to apply to the details of the specific proposals. Second, Sehon's attempt to produce a version of *AT* that distinguishes cases of action from cases like Norm's is unsuccessful.

Sehon seemingly intends to tackle Davidson's challenge head on. He writes: "Briefly put, teleological explanations support certain counterfactual conditionals, and this will allow us to distinguish the reason for which an agent acted from other nonmotivating reasons" (2005, p. 157). He contends that "when we are determining which of" two alternative teleological explanations "is true, we can look at a variety of counterfactual situations" (p. 158). "Basically speaking, we look at the agent's behavior in

the counterfactual situations and determine the goal or goals for which her behavior would have been appropriate.” “The general point,” Schon reports, “is that we are looking at counterfactual situations to see what account of the agent’s behavior makes the most rational sense. Thus, the sort of case that Davidson proposes is not enough to undermine the teleological alternative to causalism.”

A fatal flaw in Schon’s reply to the challenge is easily identified. Suppose you know Al pretty well and you know that he mowed his lawn this morning. Al’s friend Ann tells you that he had the two reasons for doing this that I mentioned, and she voices her confidence that he did it for only one of these reasons. She promises to give you \$10 if you figure out for which of the two reasons he mowed his lawn this morning and tell her how you figured it out. You decide to follow Schon’s lead and to consider various counterfactual scenarios. You know that Al dislikes mowing his lawn in even a light rain, and you start by asking yourself what he would have done this morning if there had been a light rain. You think that if he would have mowed his lawn anyway, “that is good evidence that in the actual circumstances [he] was directing [his] behavior” (Schon 2005, p. 158) at getting revenge, because the rain, for Al, would outweigh schedule-related convenience. “Would he have mowed it anyway?” you ask yourself. And you find that you are stumped. You realize that if you had substantial grounds for believing that Al mowed his lawn to get revenge, you could use those grounds to support the claim that he would have mowed it even in a light rain; and you realize that if you had substantial grounds for believing that Al mowed his lawn only for reasons of convenience, you could use them to support the claim that he would not have mowed it if it had been raining. It dawns on you that the strategy of trying to identify the reason for which Al actually acted by trying to figure out what he would have done in the counterfactual scenario I mentioned and other such scenarios puts the cart before the horse. Asking your counterfactual question about the rain scenario is nothing more than a heuristic device—and not a very useful one. The truth about what Al would have done in a light rain is grounded partly in the truth about the reason for which he actually acted.

As I have already observed in response to an earlier proposal by Schon that featured counterfactuals, the truth of true counterfactuals is grounded in facts about the actual world; and if, for example, relevant counterfactuals about Al are true for the reasons one expects them to be, their truth is grounded partly in Al’s acting for the reason for which he

acted. As far as Davidson's challenge is concerned, we are back to square one. Certain counterfactuals about Al are true partly because he acted for a certain reason. But in virtue of what is it true that he acted for that reason? Schon's proposal about counterfactuals leaves this question unanswered.

Schon has not undermined my objection to Wilson's proposal, has not produced a version of *AT* that distinguishes cases of action from cases like Norm's, and has not offered an adequate reply to Davidson's challenge. This is bad news for at least one version of *AT*.

I close this subsection with a discussion of Giuseppina D'Oro's reply to Davidson's challenge. She writes:

Since the kind of explanations which must be employed in order to do justice to the concept of actions are rationalisations, if one is to remain true to the concept of action, then arbitrating between one interpretation and another must remain a matter of choosing the interpretation which makes most sense. There is simply no way of individuating which reasons are 'efficacious' other than by asking what reasons can be invoked to rationalise the action. (2007, pp. 19–20)

What is to be done, then, when two competing "interpretations" of an action make equal sense or when neither one makes more sense than the other? D'Oro writes:

It certainly cannot be denied that there will be cases where more than one compelling rationalisation for one and the same action might be available. But even in cases where there are multiple reasons in the light of which one might render an action intelligible, the question "what is the correct interpretation?" cannot be settled by introducing the notion of a psychological process since, in the last analysis, what is crucial to the concept of action explanation is not the idea of descriptive adequacy but that of intelligibility. (p. 20)

Now, a causalist who grants that intelligibility is crucial to the concept of action explanation, as Davidson does, also maintains that an interpretation of an agent's action that represents the action as intelligible fails to be an adequate explanation of that action if it cites no cause of the action. The causalist's claim, in part, is that attractive interpretations of actions may fall short of being adequate explanations.

Why should this causalist idea be rejected? The passage from which I have been quoting continues as follows: “The fact that there may not always be clear answers in a particular case provides no grounds for conflating the conceptual question ‘what does it mean to explain something as an action?’ with the epistemological question: ‘how do we know whether the agent really acted for this reason?’” (D’Oro 2007, p. 20). But the causalist position I have been discussing does not conflate these two questions. The causalist claims that whether or not we know what reason the agent acted for, an interpretation of an action is not an adequate explanation if it cites no cause of the action. Of course, the causalist’s claim is more specific than this. Here is *D₅* again: Necessarily, if *E* is an adequate explanation of an intentional action *A* performed by an individual agent *S*, then *E* cites (1) a reason that was a cause of *A* or (2) a belief, desire, or intention that was a cause of *A* or (3) a neural realizer of a belief, desire, or intention, which neural realizer was a cause of *A* or (4) a fact about something the agent believed, desired, or intended, which fact was a cause of *A*. This is a claim about the nature of action explanation; it is not a claim about an epistemological question.

Imagine a scene in a novel in which a man named Al mows his lawn at the crack of dawn. The author, Amber, plays up two plausible motives for the early mowing, one having to do with vengeance and the other with convenience. In the novel, Al’s teenaged children have a discussion about his mowing, as do the people in the house next door. Both groups narrow Al’s likely motives down to two, and neither group comes to an agreement about which motive Al was acting from. Amber announces that one motive or the other was at work, but not both. Her aim is to move her readers to think interpretively about Al’s conduct, and she herself makes no decision about which motive Al acted from. She finds it amusing to leave this open in the fictional world she sketches and in her own mind.

Apparently, on a view like D’Oro’s, two competing interpretations of Al’s early mowing are adequate explanations of it. However, causalists will claim that because there is no fact of the matter in the novel about which motive Al was acting from, neither interpretation is an adequate explanation of the action at issue. Causalists contend that if Al were intentionally mowing in the actual world, rather than in this incomplete fictional world, there would be a fact of the matter about the motive from which he was mowing even if no one—including Al—knows what that motive is. They do not confuse a conceptual question about the nature of action explanation with an epistemic question. Do anticausalist interpretationists treat human agents more like inhabitants of incomplete fictional worlds than like flesh and blood parts

of the actual world? I leave it to readers to reflect on this question and on why I closed this subsection with this paragraph.

3.4.3. *Ginet 1990 and Wallace 1999*

Carl Ginet defends a position on “reasons explanation” that holds out the promise of answering Davidson’s challenge (1990, chap. 6). This subsection is a critique of Ginet’s position and a related idea advocated by R. Jay Wallace (1999).¹⁰

Consider a reasons explanation of the form “*S V*-ed in order (thereby) to *U*” (Ginet 1990, p. 137). “The only thing *required* for the truth of a reasons explanation of this sort,” Ginet writes, “besides the occurrence of the explained action, is that the action have been *accompanied* by an intention with the right sort of content” (p. 138).¹¹ In particular, it is not required, in his view, that the intention figure in the causation of *V* or any part of *V*. “Given that *S* did *V*,” Ginet contends, it is sufficient “for the truth of ‘*S V*-ed in order to *U*’ that “concurrently with her action of *V*-ing, *S* intended by *that* action to *U* (*S* intended *of that* action that by it she would *U*).” He adds: “If from its inception *S* intended of her action of opening the window that by performing it she would let in fresh air (from its inception she had the intention that she could express with the sentence ‘I am undertaking this opening of the window in order to let in fresh air’), then ipso facto it was her purpose in that action to let in fresh air; she did it in order to let in fresh air.”

Ginet writes: “The content of the intention is . . . the proposition ‘By this *V*-ing (of which I am now aware) I shall *U*’. It is owing to this direct reference that the intention is about, and thus explanatory of, *that particular* action” (p. 139). He asserts that an intention of this kind, because it is an intention about a particular action, “could not begin before the particular action does” (p. 139). However, he says, the action does not need to be complete before one can have an intention about it: “It is enough if the particular [action] has begun to exist.” So imagine that, for some reason or other, Ann gets up to open a window and then, while opening it, acquires an intention, *N*, of her opening it, that by so doing she let in fresh air. Would she have opened the window—performed that action—even if she had not acquired *N*? Perhaps. Maybe she set out to open the window to get a better view of the street, and perhaps she would have opened it (simply for that purpose), if she had not acquired *N* or any other intention concerning her letting fresh air into the room. Even then—that is, even if the counterfactual is true—*N* might properly figure in an explanation of Ann’s opening the window. For one thing, the

completion of that action might be causally overdetermined by *N* and Ann's intention to get a better view. *N* would not causally explain the entire action (from the beginning on), since *N* was not present until the action was already in progress. But it might enter into a causal explanation of the action's being completed.

To shine the spotlight on an important issue, I invite the reader to consider a story that would be question-begging if it were not employed simply for that purpose. God, who is omniscient and never lies, tells us that Ann had the following two *de re* intentions while opening the window and that both were present at the time of the completion of that action: Ann had the intention, *N*, of her opening the window, "that by it she would" let in some fresh air; and she had the intention, *O*, of her opening the window, that by it she would gain a better view. God also tells us that exactly one of these intentions is explanatory of Ann's opening the window while refusing to say which. He wants us to ask for additional facts in an attempt to figure out why Ann opened the window.

If you, dear reader, are open to the idea that this scenario is possible—the idea that it involves no contradiction—you may be inclined to look for some difference between *N* and *O* that might help to account for the fact that only one of these intentions helps to explain Ann's opening the window. To make a potentially long story short, suppose that you eventually ask God how Ann's relevant bodily motions at the time were produced. He replies that *O* (or its neural realizer) helps to produce bodily movements involved in Ann's opening the window and *N* (like its neural realizer) plays no causal role at all in the production of any part of the action. People who believe God's various assertions in this story would reasonably judge that *O* is the explanatory intention and *N* is just along for the ride. Such people were looking for a difference between *N* and *O* that might account for the fact that only one of the intentions is explanatory of Ann's opening the window, and the difference just identified has a strong ring of relevance.¹²

Remove God from the story and suppose that a neuroscientist, without altering the neural realization of *N* itself, renders that realization incapable of having any effect on Ann's bodily movements (and any effect on what else Ann intends) while allowing the neural realization of *O* to figure normally in the production of movements involved in her opening the window. Here, it seems, *O* helps to explain Ann's opening the window and *N* does not.¹³ Indeed, *N* seems entirely irrelevant to *why* that action was performed (a point to which I return later in this section). And if that is right, Ginet is wrong. For, on his view, the *mere presence* in the agent of an intention about her *V*-ing

(where *V*-ing is an action) is sufficient for that intention's being explanatory of her action.

Can Ginet plausibly retreat to the following position? If (1) "concurrently with her action of *V*-ing, [Ann] intended by *that* action to *U*" and (2) Ann had at the time no other intention or desire that helped to explain, either in whole or in part, her *V*-ing, then (3) the intention just mentioned is explanatory of her *V*-ing even if it (like its neural realizer) is, at the time, incapable of playing a causal role in the production of (any part of) that action. This conditional assertion will strike some readers as a nonstarter. Even if a *V*-ing is an *unintentional* action, readers may claim, some associated intentional action will have been caused in part by a desire or intention, and the desire or intention will help to explain *V*'s occurrence. For example, if, when opening a window, I unknowingly let in a fly, one might claim that, say, an intention to open the window, partly in virtue of its figuring suitably in a causal explanation of my opening the window, also figures in a causal explanation of my unintentional action of letting in the fly.

Ginet will have none of this, however. He argues that agents sometimes act in the absence of any relevant desire or intention. Some volitions, he claims, are cases in point, as are some associated "exertions" of the body. For example, he contends that "a voluntary exertion could occur [owing to an associated volition] quite spontaneously, without being preceded or accompanied by any distinct state of desiring or intending even to try . . . to exert, and it would still be an action, a purely spontaneous one" (1990, p. 9). In the case of a *voluntary* exertion of the body, Ginet says, "clearly, a causal connection between the willing and the body's exertion is required" (p. 39). But the volition itself, for Ginet, does not need to be caused (even in part) by, or concurrent with, any desire or intention.

So suppose that in Ann, standing within arm's reach of a window, a steady stream of volitions spontaneously springs up (volitions being momentary actions; Ginet 1990, pp. 32–33), as a result of which Ann's body moves in such a way as to come into contact with the window and smoothly open it in a conventional way. Suppose also that all this happens in the absence of any relevant intention or desire. Since the volitions make a causal contribution to the bodily movements that in turn cause the window to open, we have the makings of a causal explanation of all but the volitional element in Ann's opening the window. (The first and spontaneous volition in the stream is the "initial part or stage" of the voluntary exertion and the action [p. 30].) And the volitional element, on Ginet's view needs no explanation at all.

Augment the scenario by supposing that Ann intends of her opening the window “that by it she” will let in some fresh air, but that this intention, *N* (like its neural realizer), is incapable of contributing causally to the production of the bodily movements or members of the volitional stream. I do not see how *N* can have any more explanatory significance in this case than it did in the Godless two-intention case. One may think that the intention is explanatory of the action, on the grounds that, in the absence of *any* relevant intention or desire, Ann’s opening the window—that action—would be incomprehensible. But if Ginet is right, such an action requires no intention or desire at all for its occurrence: a spontaneous stream of volitions can do the work. Moreover, for those who think it bizarre that, in the absence of any relevant intention or desire, a steady stream of volitions of a kind suitable for window-opening bodily movements would occur in an agent, and who therefore want to bring some intention or desire into the explanatory picture, an extremely attractive candidate, for reasons identified earlier, is an intention that is *causally explanatory* of the supposed causally effective volitions.

I have been writing as though even if no intention is explanatory *S*’s *V*-ing, Ginet may be entitled to hold that *S*’s *V*-ing is an action. This issue merits attention. My discussion of it will help set the stage for a return to the Godless two-intention story. Ginet offers the following alleged counterexample to “the view that *S*’s *V*-ing at *t* was an action if and only if it was caused in the right sort of way by desire or intention” (1990, p. 9). Sal, who “is convinced that her arm is paralyzed,” tries “to exert her arm just in order to see what it is like to will an action ineffectually,” without “intending or wanting any such exertion actually to occur, perhaps even while wanting it very much *not* to occur.” And she succeeds in exerting her arm, an action.

One problem is that this alleged counterexample does not preclude Sal’s having a relevant desire or intention that *is* a suitable cause of her action. The candidates for relevant desires or intentions are not exhausted by the ones Ginet identifies. For example, Sal might want to discover “what it is like to will an action ineffectually,” believe that she can do this now if she wills to exert her arm, and, consequently, intend to will that. This may explain whatever effort Sal makes, thereby helping to explain, in conjunction with “the motor-neural connections to her arm . . . actually [being] in normal working order,” her exerting her arm. In the *absence* of such motivation, why *would* Sal, given Ginet’s description of the case, try “to exert her arm”? Indeed, this motivation is implicit in the story. After all, Sal tries “to exert her arm . . . *in order to* see what it is like to will an action ineffectually.”

Ginet claims, in the next paragraph, that “a voluntary exertion . . . could be caused by external stimulation of the brain” in the absence of any relevant intention or desire. Later, he asserts that someone might cause an agent to *V* voluntarily by sending signals to “the volitional part of her brain” (1990, p. 144). On Ginet’s view, the right sort of causal connection between volition and exertion is sufficient for voluntary exertion (pp. 39–44). So if, by signaling Sam’s brain, I cause a volition in him, which volition issues appropriately in some exertion of his body, he has exerted his body, and that is a voluntary action of his.

Can I cause a volition in an agent without either causing him to have a relevant desire or intention? Ginet says that volitions are not themselves desires or intentions (1990, p. 32). But this itself does not answer the question: something that is not a desire might be inextricably linked to a desire. An informed answer requires additional information about volition. Volition, for Ginet, is trying (pp. 31–32). What Ginet needs is an argument that one can try to *A* without the trying’s having any cause of a sort favored by causalists. His example of the agent who thinks her arm is paralyzed is supposed to give him what he needs, but the supposition that she lacks certain desires or intentions (for example, an intention to move her arm) does not entail that no other relevant intention (for example, an intention to find out what it would feel like to try to move her arm now) is present and plays a causal role in the production of the trying (and of her moving her arm). Similarly, Ginet has not shown that if I cause an agent to try to *A* by electronically stimulating his brain, that process does not essentially involve my electronically producing a pertinent intention that is a cause of the trying.

Return to the Godless two-intention story about Ann. It may be claimed that my discussion of that story begs the question against Ginet. He can stand fast and assert that Ann opens the window both in order to let in fresh air and in order to get a better view of the street, even though only one of those intentions (or its neural realizer) plays a causal role in the production of her opening the window. This issue merits attention. If I have been preaching to the choir (or the choir and a neutral audience), it is time to preach to Ginet.

In section 3.4.1, I presented a counterexample to Wilson’s proposed sufficient condition for an agent’s sentiently directing movements of his at a particular goal. That story featured intelligent manipulators. In a related case presented in Mele 1992a (with a different thesis as a target), mindless, attitude-insensitive forces are at work instead (pp. 248–49). A man—I will call him George—starts climbing a ladder while being undecided about which of the items he left on the roof to retrieve this time. After climbing

a few rungs, he decides to retrieve the bucket of bricks that he left up there. Once he makes that decision, his body moves as it does only because random *Q* signals from outer space provide exactly the right input to his muscles and, even so, it seems to George that he is in fact moving himself up the ladder in just the way that he had been doing. Coincidentally, the *Q* signals strike just as bizarre *Z* rays from Venus prevent events in his brain from causing muscle contractions, and the like. George intends of his movements that they result in his getting the bricks. I wrote:

Even though there is a reading of “The man went up the ladder” on which the sentence is true, it is *false* that “he went [the rest of the way] up the ladder because he wanted to retrieve [the bricks],” and false as well that “he went [the rest of the way] up the ladder because he (thereby) intended to satisfy his desire to get [the bricks].” Rather, he went the rest of the way up the ladder because the *Q* signals provided such and such input to his muscles. (Mele, p. 249; embedded quotations are from Wilson 1989, p. 288)

Note that I make three claims here involving the word “because,” two alleging falsehoods and one alleging a truth.

This story does not falsify Ginet’s position. The *de re* intentions featured in his view are about actions, and George’s trip up the ladder is not an action (nor a collection thereof). However, following the lead of Randolph Clarke (2010, pp. 29–30), one can tell a story in which a manipulator uses a chip that he has installed in George’s brain to cause volitions that issue in George’s climbing the ladder. The manipulator in no way wishes to assist George in getting the bricks; in fact, his plan is to cause George to retrieve his toolbox instead when he gets to the roof. About his own case, featuring arm-raising, Clarke writes: “She does not raise it because she wants to acquire the painting; she raises it because [the manipulator] causes her to raise it” (p. 30). This echoes the “because” claims in the passage from Mele 1992a reproduced in the preceding paragraph. And in this new version of my ladder story, George’s *de re* intention, by Ginet’s own lights, is about his climbing the ladder—an action. George intends of his climbing the ladder that it result in his getting the bricks.

It may be claimed that Clarke’s “because” claim begs the question against Ginet (see Ginet 2008, p. 231). If it does, then so does the following claim: (*BG*) George does not climb the ladder because he wants (or intends) to get the bricks, nor because he intends of his climbing that it result in his getting the

bricks, and he does not climb it in order to get the bricks; instead, he climbs it because of what the manipulator did to him. Is *BG* question begging?

Someone may contend, on the following grounds, that George does climb the ladder in order to get the bricks. Getting the bricks was the purpose he had in mind for climbing the ladder while he climbed; he climbed it while having an intention, of his so doing, that it result in his getting the bricks. Suppose there is a reading of “*SA*-ed in order to *B*” according to which the truth of these claims about purpose and intention is sufficient for its being true that George climbed the ladder in order to get the bricks. On this reading, the true assertion that George climbed the ladder in order to get the bricks does not yield an adequate *explanation* of his climbing the ladder. His bricks-involving desires, intentions, and reasons are no more explanatory of his climbing the ladder than they are of his nonactional trip up the ladder in the version of the story featuring mindless forces. Someone who regards *BG* as question begging may find the idea that George climbed the ladder in order to get the bricks appealing without recognizing that what can be said in favor of it does not support the crucial claim at issue—namely, that there is an acceptable noncausal *explanation* (in terms of reasons) of his action.

The point just made merits emphasis. One may distinguish between a weaker and a stronger reading of “*SA*-ed in order to *B*.” On the weaker reading, the following fact is sufficient for its being true that *SA*-ed in order to *B*: *B*-ing was a purpose he had in mind for *A*-ing while he *A*-ed; and he *A*-ed while having an intention, of his so doing, that it put him in a position to *B* or bring about his *B*-ing. On the stronger reading, a necessary condition of the truth of any statement of the form “*SA*-ed in order to *B*” is that it provides an explanation of *A*. If the weaker reading were at work in *BG*, *BG* would be question-begging. I offered no support at all for the claim that the weaker reading is not satisfied in George’s case. But it is the stronger reading that is in play in *BG*. And why is that? Because the view of Ginet’s under consideration is explicitly a view about reasons *explanations* of actions.

In the Godless two-intention story about Ann and the window, the claim that she opened it in order to let in fresh air is true on the weaker reading even though her intention regarding fresh air played no causal role in the production of her window-opening action. But it certainly does not follow from this that the claim is true on the stronger reading. Moreover, her intention regarding fresh air is no more explanatory of her opening the window than George’s intention regarding bricks is of his climbing the ladder. *Why* did George climb the ladder? Because of what the manipulator did to him, and not because of any intention George had.¹⁴

Earlier, I said that I would challenge two claims about cases in which agents who intentionally *A* have two or more reasons for *A*-ing: the claim that, in all such cases, the agents *A* for all of these reasons; and the claim that, in all such cases, there is no fact of the matter about which are the reasons for which the agents *A*. I had my discussion of the Godless two-intention story in mind. Ann had at least two reasons for opening the window—reasons that can be inferred from her intentions. If I am right, she opened it for one of those reasons and not the other.

Ginet has not persuasively answered the Davidsonian challenge. I argue next that Wallace's related attempt also fails. Wallace's response resembles Ginet's in featuring intentions. He contends that agents' intentions "incorporate information about [their] conception of their reasons for acting as they do" (1999, p. 240; also see McCann 1998, chap. 8). For example, Al, in a scenario I sketched earlier, might intend *to mow his lawn this morning as a way of getting back at his neighbor*, where the italicized words are an expression of the content of his intention.¹⁵ On Wallace's view, the reason for which Al mows his lawn is "reflected in the content of [his] intention," and agents are "guided by their conception of their reasons when that conception is reflected in the content of the intention on which they act" (1999, p. 239). While leaving it open that decisions and intentions are causes of actions, Wallace rejects the idea that intentions and decisions have beliefs and desires as causes (p. 241, n. 35). Thus, that Al mows his lawn for a reason having to do with getting back at his neighbor can be read off from an intention that plays a suitable causal role in producing the relevant bodily motions even though his desire for revenge and his belief that mowing his lawn this morning would serve that purpose (and their neural realizers) play no causal role in the production of the intention or the action.

This proposal pushes the issue back a step. Wallace and I agree that an agent's deciding to *A* is itself an intentional action (Wallace 1999, pp. 236–37). So he should see Davidson's challenge as applying straightforwardly to *deciding* for reasons. Recall that Al has a pair of reasons for mowing his lawn this morning, one centrally involving revenge (*R₁*) and the other convenience (*R₂*), but he mows it for one and not the other. Suppose that Al decides to mow his lawn this morning, but leave it open that this description of what he decides is incomplete. If it was for reason *R₁*—or reason *R₂*—that he made his decision, in virtue of what is that true? Now, it is plausible that in ordinary cases of executing a decision to *A*, or executing the intention to *A* formed in so deciding, the reasons for which we *A* are the reasons for which we decided as we did.¹⁶ So the answer to my question, on a view like Wallace's, may be that

it is true that Al decided for certain reasons—the same reasons for which he acted, reasons that can be read off from the content of his decision—in virtue of the content of his decision (that is, the content of the intention he formed in making his decision). For example, if what Al decided was to mow his lawn as a way of getting revenge, then he decided for reason *R_I*, and that in virtue of which it is true that he decided for *R_I* is precisely that *what he decided* was to mow his lawn as a way of getting revenge.

This answer is problematic. It is implausible that it is a *general truth* about our decisions to *A* that the reasons for which we so decide—and the reasons for which we act when we execute decisions—are expressed or “reflected” in the contents of our decisions. We consider many reasons for and against accepting certain job offers, for example, and sometimes we reach the decisions we do in these cases—and accept or reject a job offer—for a whole raft of reasons. It is unlikely that large rafts of reasons can be read off from the contents of our decisions in such cases. It would take a very special mind to represent each member of a large collection of reasons in the content of a decision. If, in cases of this kind, people should say (as, in fact, they do say) and believe that what the agent decided was to accept job offer *X*—and not that what he decided was, for example, to accept *X* “as a way of” bringing it about that he and his family live in a more attractive part of the world, enabling his children to attend better schools inexpensively, improving his family’s job prospects, reducing his teaching load, increasing his salary, and so on—special grounds need to be offered for holding that, even in relatively simple cases, (partial) representations or “reflections” of each of the reasons for which agents decide and act as they do uniformly enter into the contents of decisions.

There is a related problem. Suppose that Al decided for reason *R_I*. On one view, what he decided was to mow his lawn early this morning (“to *M*,” for short). On Wallace’s alternative view, what he decided was (at least) to mow his lawn early this morning as a way of getting back at his neighbor (“to *M*^{*}”). If Al decided to *M*^{*}, for what reason did he so decide? If the answer is “no reason,” then, unless Wallace is prepared to defend the thesis that there are intentional actions that are done for no reason and that some decisions are among them, he should retract his claim that decisions are intentional actions.¹⁷ (In my view, the retraction would be a mistake; see chapter 2.) So suppose there was a reason for which Al decided to *M*^{*}. On Wallace’s view, apparently, that reason is reflected in the content of Al’s decision. Now, that mowing his lawn early this morning would be a way of getting revenge on his neighbor—or, on another view of reasons, the combination of Al’s desire for revenge and his belief that he can get it by mowing early—is a reason for *M*-ing and a

reason for deciding to M , not a reason for M^* -ing and for deciding to M^* .¹⁸ The answer to my question about Al's reason for deciding to M^* is not found in *this* reason. If a positive answer is forthcoming, the reason identified seemingly needs to be reflected in the content of Al's decision, on Wallace's view. So what I had been describing as Al's decision to M^* is really a decision to M^{**} —say, a decision to mow his lawn early as a way of getting revenge on his neighbor, partly just for the sake of getting revenge, but also both in order to show her that he is not the sort to take rude mowing behavior lying down and to honor a tit-for-tat principle of his. Of course, if it is claimed that this is what Al really decided to do, the question arises for what reason he decided to do this. A vicious regress threatens and must somehow be blocked.

Wallace asserts that “we do not for a minute need to think that it is necessarily a simple matter, even for agents themselves, to ascertain what their real intentions in acting are” (p. 240).¹⁹ If the content of an intention like Al's when he mows his lawn were as complex as Wallace is apparently committed to viewing it as being, we should not be at all surprised about agents' difficulties in this connection! Worries about self-deception are another matter entirely. Even if Wallace can block the threatened regress, his view has the consequence that the contents of decisions are implausibly complicated even in mundane scenarios like the present one.

I supposed that Wallace would not want to deny that Al's decision to M^* was made for a reason. However, that supposition is not required for my purposes. In the story as I sketched it, if Al decided to mow his lawn early as a way of getting back at his neighbor, he made this decision for the reasons I identified. If Wallace were to deny that Al's decision to M^* was made for a reason, he would be wrong.

Here is the bottom line on Wallace's reply to the Davidsonian challenge. Wallace does not answer the challenge as it applies to mental actions of intention formation—that is, decisions. A natural answer on his behalf, given his position on acting for reasons, is unsuccessful. The contents of our decisions and intentions are not equipped to do the required work. Nor can they do the required work regarding overt actions done for complex collections of reasons. In simple cases, one may think that it is true that the reasons for which an agent acted were R in virtue of R 's being reflected in the content of his decision, even though one denies that this is true in many cases. But this stance is unstable. Actual psychological constraints on the complexity of the content of a normal human agent's decisions may permit contents that reflect the reasons, R , for which an agent decided and acted in some cases, but that an effective decision had R -reflecting content certainly does not

entail that that in virtue of which it is true that the agent decided and acted for *R* is that the decision had that content. A credible *general* answer to the “in virtue of” question, one that works in all cases of acting (including deciding) for reasons, is what theorists are after.²⁰ And causalism has resources for providing such an answer. Perhaps in relatively simple cases (for example, a father’s deciding to try to cheer his daughter up, which he intrinsically desires to do, by throwing a party for her), the reason for which the agent decided and acted as he did—reason *R*—can be read off from the content of his decision. This, of course, is entirely consistent with its being true that he decided and acted for *R* in virtue of its being true that *R* (or his having or apprehending *R*, or the neural realization of one of these things, or some fact about his relation to *R*) played a distinctive causal role in generating the decision and overt action. The latter truth, in conjunction with the supposition that *R* was reflected in the content of the agent’s decision, would account for its being true that the reason reflected there was the reason for which he decided and acted.

3.5. Causalism vs. Anticausalist Teleologism

What can causalists about action explanation and proponents of *AT* (anticausalist teleologism) reasonably challenge one another to do? Distinguishing the project of producing a *conceptually sufficient condition* for a human being’s acting in pursuit of a particular goal from the project of producing an *analysis* of this is useful in answering this question. The latter project obviously is more challenging than the former: producing conditions that are individually conceptually necessary and jointly conceptually sufficient for *X* is more challenging than producing a conceptually sufficient condition for *X*. And if it were demonstrated, for example, that having a mental state or event of a certain kind among its causes is a *necessary* condition for an event’s being an action performed in pursuit of a goal *G* and that no attempted explanation of an action that did not appeal either explicitly or implicitly to such a cause is an adequate explanation, that demonstration alone would show that *AT* is false. A causalist does not need to generate an acceptable analysis of an action’s being performed in pursuit of *G* in order to show that *AT* is false.²¹

Similarly, a proponent of *AT* does not need to generate an acceptable analysis of an action’s being performed in pursuit of *G* in order to show that causalism about acting in pursuit of a goal is false. I have never challenged proponents of *AT* to produce an *analysis* of this. (To ask for an analysis of any contested philosophical concept that will be widely accepted is to ask for a

great deal.) My challenge to proponents of *AT* concerns a much more modest task: producing an informative conceptually sufficient noncausal condition for an action's being performed in pursuit of a goal *G*. Teleologists contend that "teleological explanations explain by specifying an action's goal or purpose; for example, when we say that Jackie went to the kitchen in order to get a glass of wine, we thereby specify the state of affairs at which her action was directed" (Sehon 1997, p. 195). But in virtue of what is it true that a person acted in pursuit of a particular goal? Despite their considerable efforts, Wilson, Sehon, D'Oro, Ginet, and Wallace have failed to answer this question successfully, as I have argued here.

Can causalists do any better? Elsewhere, I have argued that they can. In chapter 2 of Mele 2003a, I offered, from a particular causalist perspective, what I argued to be an informative conceptually sufficient condition for a being's acting in pursuit of a particular goal. And I observed that if the condition I offered is indeed conceptually sufficient for this, causalists are in much better shape in this connection than proponents of *AT*, none of whom have succeeded in offering noncausal conceptually sufficient conditions for this (p. 59). Interested readers are encouraged to consult Mele 2003a, chap. 2. The present chapter is concerned with related business.

If I had ever thought I knew of a convincing *direct* argument for the thesis that an agent's acting in pursuit of a particular goal requires the satisfaction of a specific causal condition, I would have gone public with it. My way of defending causalism has been indirect. I have challenged proponents of *AT* to produce informative conceptually sufficient noncausal conditions for acting in pursuit of a particular goal, argued against their proposals, and developed a causalist view of what actions are, how they are explained, and how they are produced (Mele 1992a, 2003a). If readers who compare my causalist view with anticausalist alternatives in light of the arguments I offered for the former and against the latter are justified in judging that my view is considerably closer to the truth, I am satisfied.²²

Notes

1. An objective favorers theorist about reasons for action may ask how reasons are involved in explanations of intentional actions that are done for reasons. Here is a short answer: If D_3 is true, reasons are involved in explanations of such intentional actions in a way consistent with the truth of D_3 . The task of developing a detailed answer is left as an exercise for the reader. But see Mele 2003a, pp. 79–84, for some guidance.

2. For a justification of the instruction in brackets, see Mele 2003a, p.47.
3. This paragraph and the next eight are borrowed, with some minor modifications, from Mele 2003a, pp. 48–50.
4. See Adams and Mele 1992, p. 325; Armstrong 1980, p. 71; McCann 1975, pp. 425–27, and McGinn 1982, pp. 86–87.
5. See James 1981, pp. 1101–3. For discussion of a case of this kind, see Adams and Mele 1992, pp. 324–31.
6. Readers who regard the claim that Norma brought it about that she got across the room as an exaggeration may be happy to grant that she helped to bring that about. Her helping to do that is an action.
7. On a case that may seem to be problematic for T_3 , see Mele 2003a, p. 64–65, n. 22.
8. The remainder of this subsection derives from Mele 2003a, pp. 50–51.
9. In my original story, the Martians interfere with Norm only “on rare occasions” (Mele 2003a, p. 49). So, seemingly, they may interfere with him much less often than Sally has her finger problem. The variant of Norm’s case just sketched renders speculation about this comparative issue otiose.
10. This subsection derives largely from Mele 2003a, pp. 39–42, some of which derives in turn from the discussion of Ginet’s position in Mele 1992a, pp. 250–55.
11. Ginet takes for granted the existence of intentions as states of mind, and I follow suit in my discussion of his position. Readers who deny that there are states of mind will view Ginet’s position (as he formulates it) as a nonstarter.
12. It may be objected that an omniscient, honest God could not possibly say this, since Ann’s alleged intentions are actually a single, complex intention—an intention, of her opening the window, that by it she will gain a better view and let in some fresh air. However, it is not conceptually necessary that the two alleged intentions agglomerate, even if such intentions normally do agglomerate; and a conceptual possibility is enough for present purposes.
13. It may be objected that the human brain is such that the identified effect of the scientist’s tinkering would require a change in the neural realization of N itself. Even if this is so, Ginet’s analysis of “ S V -ed in order (thereby) to U ” is a perfectly general conceptual analysis, and it is conceptually possible that an agent have a brain that allows the scientist to accomplish his trick.
14. Readers who reject Ginet’s claim that there can be volitions in the absence of any relevant intention will reject an assumption in my story about manipulated George. They will say that the manipulator can cause the volitions at issue only if he causes George to have a relevant intention. Let it be an intention to get his tool kit. On his way up the ladder, George acquires an intention, of his climbing, that it result in his getting the bricks. The latter intention has no effect at all on George’s climb. Did George, even so, climb the ladder in order to get the bricks? In the weaker sense, yes. But I see nothing to recommend the claim that he also does so in the stronger sense.

15. In this formulation of the content of Al's alleged intention, I follow Wallace. He writes: "*A*'s intention is to provide assistance as a way of doing what is right, while *B* acts on the different intention of providing assistance as a way of collecting a financial reward" (1999, p. 240).
16. On other cases, see Mele 1992b and 1995b.
17. For an exceptional case in which an agent does something intentionally but not for a reason, see Mele 1992b.
18. Discussion of this issue is complicated by my neutrality on action-individuation (see chapter 1, section 1.2). On a fine-grained view, Al's mowing his lawn early (*M*) and his mowing his lawn early as a way of getting back at his neighbor (*M*^{*}) are two different actions. Any reason for which Al *M*^{*}'s encompasses something that explains his acting to get back at his neighbor, but a reason for which he mows his lawn early does not need to do this. For example, in reporting that a reason for which Al mowed his lawn early (i.e., *M*-ed) was to get back at his neighbor, one does not explain the "as a way of getting back at his neighbor" aspect of his *M*^{*}-ing. On a coarse-grained view, *M* and *M*^{*} are the same action under different descriptions and effective reasons are relativized to action-descriptions. Any effective reason for Al's action under description "*M*^{*}" encompasses something that explains his acting to get revenge on his neighbor, but an effective reason for his action under description "*M*" does not need to do so. A componential view of action-individuation yields a similar result.
19. On an alternative view, the claim would be that it is not always easy, even for the agents, to know for what *reason(s)* they are *A*-ing. For example, Al might believe that it was for reasons of convenience that he decided to mow his lawn early this morning and that he is now mowing it for those reasons, whereas, in fact, it was for reasons of vengeance that he decided to mow it and he is mowing it for the latter reasons. Again, on Wallace's view, the reasons for which Al made his decision and for which he mows his lawn can be read off from his intention, an intention that Al has without realizing it. On the alternative view, the pertinent reasons are the ones that played a suitable causal, explanatory role in the production of Al's decision to mow and his mowing, even though Al does not realize that these reasons are the operative ones.
20. Thus, my points about representational limitations, for example, obviously cannot be accommodated by claiming simply that although, in some cases, only *some* of the reasons for which an agent decided to *A* can be read off from the content of his decision, they are reasons for which he so decided in virtue of that. Of course, a general answer can be disjunctive, but a disjunctive general answer will provide all the disjuncts.
21. Incidentally, I have never offered an analysis of action nor of acting in pursuit of a particular goal, although I have defended causalism in both connections (Mele 1992a, 2003a). Paul Moser and I (Mele and Moser 1994) have offered an analysis of what it is for an action to be an intentional action.

22. The articles of mine on which parts of this chapter are based are Mele 2010 and 2013a. I am grateful to Andrei Buckareff and Randy Clarke for comments on a draft of the former and to Giuseppina D'Oro and Scott Schon for comments on a draft of the latter. (Parts of this chapter also derive from two books of mine, Mele 1992a and 2003a, as indicated in notes 3, 8, and 10.)

Agents' Abilities

CLAIMS ABOUT AGENTS' abilities—practical abilities—are common in the literature on free will, moral responsibility, moral obligation, personal autonomy, weakness of will, and related topics. These claims often ignore differences among various kinds or levels of practical ability. In this chapter, using *A* as an action variable, I distinguish among three kinds or levels: simple ability to *A*; ability to *A* intentionally; and a more reliable kind of ability to *A* associated with promising to *A*. I believe that attention to them will foster progress on the topics I mentioned. My focus is kinds, levels, or grades of practical ability—not the metaphysics of ability. My hope is that theorists who disagree with one another about the metaphysics of ability will be persuaded to agree with me that the distinctions I draw among kinds, levels, or grades of ability are useful.

4.1. Two Kinds of Specific Practical Ability

Although I have not golfed for years, I am able to golf. I am not able to golf just now, however. I am in my office now, and it is too small to house a golf course. The ability to golf that I claimed I have may be termed a *general* practical ability. It is the kind of ability to *A* that we attribute to agents even though we know they have no opportunity to *A* at the time of attribution and we have no specific occasion for their *A*-ing in mind. The ability to golf that I denied I have is a *specific* practical ability, an ability an agent has at a time to *A* then or to *A* on some specified later occasion.¹ My specific concern in this chapter is specific abilities.

There is an ordinary sense of “able” according to which agents are able to do whatever they do.² In this sense of “able,” an agent’s having *A*-ed at a time is conceptually sufficient for his having been able to *A* then. If Ann backed

her car into mine, she was able to do that, in this sense. That is so whether she intentionally or accidentally backed her car into mine. Similarly, if Ann threw a basketball through a hoop from a distance of ninety feet, she was able to do that in this sense, and that is so whether she was trying to throw it through the hoop, or simply to hit the backboard, or merely to throw it as far as she could. Yesterday, Ann rolled a six with a fair die in a game of chance. She was able to do that, in the sense of “able” at issue.

I said that there is a sense of “able” in which these claims are true. It can also be said that there is a kind of ability about which claims such as these are true. I call it *simple ability*. I have not claimed that simple ability to *A* is found *only* in cases in which agents *A*. Rather, my claim is that an agent’s *A*-ing at a time is sufficient for his having the simple ability to *A* at that time. Another condition that may be sufficient for this is discussed in section 4.2.

Being simply able to *A* is distinguishable from being able to *A intentionally*. It is controversial how much control agents who *A* must have over their *A*-ing in order to *A intentionally*. Even so, there are clear illustrations of a difference between control that is appropriate for intentional action and control that falls short. Ann has enough control over her body and dice to roll a die intentionally, but, like any normal human being, she lacks control over dice needed for rolling a six intentionally with a single toss of a fair die. Therefore, although she is able to roll a six with a single toss of a fair die, she is not able to do that intentionally. Her throwing a six now owes too much to luck to be intentional. Even if, wrongly thinking that she has magical powers over dice, Ann *intends* to throw a six now and does so, she does not intentionally throw a six. A proper account of being able to *A intentionally* hinges on a proper account of *A*-ing intentionally and the control that involves. Paul Moser and I have offered an analysis of intentional action (Mele and Moser 1994), but there is no need to insist on that analysis here. However intentional action is to be analyzed, being able to *A intentionally* entails having a simple ability to *A* and the converse is false.³ Noticing that the former ability is stronger than the latter in this sense suffices for present purposes. I have no need here for an *analysis* of being able to *A intentionally* or of the control intentional action requires.

A confusion about control should be identified. Sometimes it is claimed that agents have no control at all if determinism is true. The claim is false. When Ann drives her car (under normal conditions), she controls the turns it makes even if her world is deterministic. She plainly controls her car’s movements in a way that pedestrians and her passengers do not. For example, she turns the steering wheel and they do not. A distinction can be drawn between

a kind of agential control that is compatible with determinism and a kind that is not.⁴

It will be useful to have an easy way of moving back and forth between "ability" claims and "able" claims in terms of the distinction I sketched. I abbreviate "simple ability to *A*" as "*S*-ability" and "ability to *A* intentionally" as "*I*-ability." Corresponding "able" expressions are "*S*-able" and "*I*-able."⁵

4.2. Could Have Succeeded

The topic of the present section is a pair of views: a commonsense view of agents' *S*- and *I*-abilities; and a view of agents' abilities favored by libertarians. Libertarians have the option of distinguishing between directly and indirectly free actions (Mele 1995a, pp. 207–9). For example, they can claim that *A* is a directly free action only if its proximal causes did not deterministically cause it and claim as well that an action that was deterministically caused by its proximal causes may be indirectly free, provided that the agent earlier performed some relevant directly free action or actions. A more specific example is provided by the claim that an agent might have freely pressed a button even if his pressing it was deterministically caused by proximal causes that included his deciding to press it straightaway, if his deciding to do that was a directly free action. According to libertarians who countenance indirectly free actions, only agents who have performed directly free actions can perform indirectly free actions. This gives directly free actions a special importance: without them, there are no free actions at all.

Typical libertarians maintain that performing a directly free action requires being able to perform an alternative action and that determinism precludes this ability.⁶ I will motivate the suggestion that a commonsense view of *S*- and *I*-ability might be silent on the question whether determinism precludes this. The suggestion's plausibility enables me to move forward without attempting to resolve a long-standing dispute between libertarians and traditional compatibilists.

Here are two pronouncements of common sense (*CS*, for short). First, we have both general and specific abilities to do things we never do. Although Beth was able to buy a plane ticket to Beijing, she never did. Thirty years ago, on her seventieth birthday, she was tempted, for the first time, to book a flight to Beijing and was able then to do so straightaway, but she decided against the purchase and never again considered flying there. Second, we occasionally try and fail to do things we are *S*-able to do and things we are *I*-able to

do. A skilled putter may fail to sink the next three-foot putt he attempts even though he was *S*-able and *I*-able to sink it.

Libertarians and other incompatibilists typically hold that an agent who did not *A* at *t* was able to *A* at *t* only if in another possible world with the same past and laws of nature, he *A*-s at *t*.⁷ On this view, if agents in deterministic worlds are able to do anything at all, they are able to do only what they actually do. For in any world with the same past and laws as *S*'s deterministic world, *Wd*, *S* behaves exactly as he does in *Wd*. For my purposes in this chapter, I have no stake in accepting or rejecting this view, provided that it can be understood as a view about a *species* of ability. I will suppose that there is a species of ability—*L*-ability—such that, by definition, an agent *S* in *W* has, at the relevant time, the *simple L*-ability to *A* at *t* if and only if there is a possible world with the same past and laws as *W* (either *W* itself or another world) in which *S A*-s at *t*.⁸ Similarly, I will suppose that, by definition, an agent *S* in *W* has, at the relevant time, the *L*-ability to *A intentionally* at *t* if and only if there is a possible world with the same past and laws as *W* in which *S A*-s intentionally at *t*. One virtue of these accounts is their precision.

It may be argued that any view of *S*- and *I*-ability that makes the two pronouncements I identified presupposes that determinism is false. But such an argument may expect too much of *CS* views of these abilities. Consider a superb free-throw shooter, Peta. Owing to years of practice and the skills she developed, she sinks about 90% of her free throws and typically is *I*-able to sink a free throw. Sometimes, when Peta misses, she has been fouled very hard and sees stars or is dizzy. Normally, however, things just do not go quite right when she misses. Peta may release the ball a little too early or too late, throw it a little too hard or too soft, push a bit too much or too little with her legs, or the like. If Peta's world is deterministic, all occurrences of these problems are deterministically caused. But what *CS* says about *I*-abilities may not be metaphysically deep. Perhaps, on a *CS* view of *I*-ability, that, under normal conditions, an agent intentionally *A*-s in the great majority of instances in which she attempts to *A* and that the conditions under which she just now tried to *A* were normal is sufficient for her having had the ability to *A* intentionally at the time—even if her attempt failed. If what *CS* says about *I*-ability is inseparable from its alleged claims about *freedom-level ability*, discussion of familiar issues dividing compatibilists and incompatibilists would be in order now.⁹ However, it is conceivable that a *CS* view of *I*-ability is silent on freedom-level ability, that it takes no explicit stand on whether determinism is true or false, and, indeed, that it ignores the topic of determinism.

What about simple ability—in particular, an *S*-ability to *A* possessed by an agent who is not able to *A* intentionally? When Fred tosses a fair die, he tosses a six about a sixth of the time. He has experimented with ways of trying to roll a particular number—a six, for example—but he has not developed any special dice-rolling skills. Just now, Fred is playing a board game and is about to roll a die. If his world is deterministic, then whatever number he tosses, his tossing that number is deterministically caused. Suppose he throws a five. Was he able to throw a six? Perhaps, according to *CS*, that, under normal conditions, an agent *A*-s (for example, rolls a six) about a sixth of the time he *B*-s (for example, rolls a die) and that the conditions under which he just now *B*-ed were normal is sufficient for his having been *S*-able to *A* at the time.¹⁰

The conditions that I suggested may suffice, according to *CS*, for Peta's being *I*-able to sink her free throw and for Fred's being *S*-able to toss a six on his next roll are compatible with Peta's missing her shot and with Fred's rolling a five, as in fact they did. *CS* folks should welcome this point. After all, the fact that an agent who is acting at a time is able then to *A* then, intentionally or otherwise, is not commonly regarded as entailing that he *A*-s. (I am able to dress in a kilt, but I doubt I ever will. Just now, a friend gave me a kilt and dared me to wear it to lunch. I declined.) If we know that Peta's and Fred's worlds are deterministic, then we know, given how things turned out, that the state of their worlds millions of years ago and the laws of nature are such that, at *t*, Peta misses her shot and Fred rolls a five. But "*S* does not *A* at *t*" is not commonly regarded as entailing "*S* is unable at *t* to *A* at *t*." Its being causally determined that *S* will not *A* at *t* does entail that *S* lacks certain *L*-abilities. However, conceivably, some *CS* folks who are compelled to think about determinism may judge, consistently with their *CS* view of *S*- and *I*-ability, that *L*-abilities have compatibilist analogues—that is, analogues compatible with determinism. Other *CS* folks may judge that there can be no such analogous abilities. But such a judgment may reach beyond their *CS* view of *S*- and *I*-ability rather than being an implicit pronouncement of it.

Philosophers happy to talk in terms of possible worlds will say that an agent in a world *W* is *S*-able to *A* at *t* if and only if he (or a counterpart) *A*-s at *t* in some relevant possible world, and is *I*-able to *A* at *t* if and only if he (or a counterpart) *A*-s intentionally at *t* in some relevant possible world.¹¹ One way to see the disagreement between incompatibilists and compatibilists about determinism and being able to do otherwise is as a disagreement about what worlds are relevant. According to incompatibilists about this, all and only worlds with the same past and natural laws as *W* are relevant; they hold the past and the laws fixed. Compatibilists disagree. I have been suggesting, in

effect, that a representative of *CS* who is forced to think about possible worlds may take the following position on an agent who, at *t*, tried and failed to *A* or played a game of chance in which he “took a chance” at *A*-ing but did not *A*, *A* in both cases being a kind of action the agent has often performed: relevant worlds include all worlds with a very similar past and natural laws in which the agent has the same “*A*-rate” under normal conditions (for example, the same rate of intentionally sinking a free throw, and the same rate of throwing a six when he throws a fair die) and in which conditions are normal at the relevant time. This is not to say, of course, that these are the *only* worlds that may be deemed relevant. After all, one would want to leave room for agents’ being able to do things in abnormal circumstances. For example, Peta presumably is able to sink a free throw, and may be able to sink it intentionally, even when the hoop’s circumference is slightly smaller than normal. One would also want to leave room for abilities to do “new things”—for example, sinking one’s first free throw or putt.

The expression “normal conditions” cries out for attention. *CS* may not say anything very detailed about it, however. Presumably, normal conditions in Fred’s case exclude such things as a lopsided die and a properly shaped and weighted die that is being controlled by fancy machines. But perhaps as *CS* understands normal conditions, they do not exclude a combination of normal gravitational forces and normal velocities, trajectories, spins, and bounces of normal dice that may be a major part of a deterministic cause of Fred’s die’s landing five-up. Similar points may be made about Peta’s case. Normal conditions exclude a deformed basketball, a smaller than normal hoop, dizziness, blurred vision, and the like. However, perhaps normal conditions, as *CS* understands them, do not exclude various small-scale bodily events that are in the normal range for Peta when she is attempting a free throw but may add up to a major part of a deterministic cause of her shot’s hitting the rim and bouncing away.

The simplicity of the accounts I suggested of simple *L*-ability to *A* and *L*-ability to *A* intentionally is attractive. I doubt that equally simple, promising accounts of *S*-ability and *I*-ability are accessible from the commonsense perspective on these abilities that I have been discussing.¹² In the literature, what looks like the most promising account, from this perspective, of something resembling *I*-ability—an analysis of a kind of responsiveness to reasons—is intricate. Incidentally, it comes, not from traditional compatibilists, but from semicompatibilists, philosophers who hold that determinism is compatible with free action and moral responsibility even if it is incompatible with agents’ ever having been able to act otherwise than they did (Fischer 1994;

Fischer and Ravizza 1998).¹³ Semicompatibilists contend that free action and moral responsibility do not require an ability of this kind, and they do not need to be in the business of providing an analysis of being able to *A*. In any case, although simplicity has its virtues, a true appeal to greater simplicity would not show that there are not, in addition to simple *L*-ability to *A* and *L*-ability to *A* intentionally, non-*L* analogues of these abilities in some deterministic worlds.

4.3. A Third Kind of Practical Ability: Some Preliminaries

"Able to *ensure* that *p*," an expression Peter van Inwagen uses in an intriguing article (2000, p. 8), has the ring of intended reference to an ability that is more robust or reliable than a garden-variety ability intentionally to bring it about that *p*. Is there a kind of ability to *A* that is more reliable than the ability to *A* intentionally, perhaps an *ensuring* ability? In this section, I take my lead from van Inwagen in preparing to search for an ability of this kind.

Actually, van Inwagen's expression, "able to ensure that *p*," is not a happy one for what he seemingly has in mind. Let *p* be "Ann's basketball goes through Bob's hoop at *t*." Then one thing that would ensure that *p* is true is Shaq's slam-dunking Ann's ball through Bob's hoop at *t*. Another is Shaq's shooting it through the hoop at *t* from the free-throw line. Yet another is Ann's shooting it through the hoop at *t* from ninety feet away. Shaq and Ann are at least *S*-able to do these things, which things would ensure that the ball goes through the hoop at *t*. (Shaq also is *I*-able to slam-dunk the ball through the hoop.)¹⁴ They are able to ensure that the ball goes through the hoop insofar as (or in the sense that) they are able to do things that ensure that the ball goes through the hoop. This is not the sort of thing van Inwagen is after, as will soon become clear. What kind of practical ability might he have in mind?

Attention to part of an argument in van Inwagen 2000 against the theoretical utility of agent causation will help narrow the possibilities.¹⁵ A central plank in the argument is, roughly, the claim that an agent who knows that "it is undetermined" (p. 17) whether he will *A* is not able to *A*. Van Inwagen's defense of this claim features a scenario in which he knows, perhaps because God told him, that there are "exactly two possible continuations of the present, . . . in one of which" he reveals a damaging fact about a friend to the press "and in the other of which" he keeps silent about his friend (p. 17). He also knows that "the objective, 'ground-floor' probability of [his] 'telling' is 0.43 and that the objective, 'ground-floor' probability of [his] keeping silent

is 0.57.” Van Inwagen says that he does not see how he can “be in a position to” promise his friend that he will keep silent. He adds:

But if I believe that I am able to keep silent, I should, it would seem, regard myself as being in a position to make this promise. What more do I need to regard myself as being in a position to promise to do X than a belief that I am *able* to do X? Therefore, in this situation, I should not regard myself as being able to keep silent. (And I cannot see on what grounds third-person observers of my situation could dispute this first-person judgment.) (2000, pp. 17–18)

This, van Inwagen says, is an “argument for the conclusion that it is false that I am able to keep silent” (p. 18).¹⁶

To eliminate a source of distraction, I suggest that van Inwagen’s claims about promising be understood to be about *sincere* promising. Another source of distraction should also be eliminated. There may well be a difference between the probability that van Inwagen will keep silent and the probability that he will keep silent given that he promises to keep silent. I will assume that the 0.57 probability van Inwagen mentions is the probability of the latter.

There are many things I believe I am able to do that I do not “regard myself as being in a position to promise [sincerely] to do”—for example, toss heads now with the quarter I am holding. My belief that I am able to do this is an utterly ordinary belief. The kind of ability it is about is what I called simple ability. Van Inwagen’s belief that he is not able to keep silent in the imagined scenario presumably is not about simple ability. We who believe that “the objective, ‘ground-floor’ probability of [his] keeping silent is 0.57” can easily imagine that he does keep silent. If he keeps silent, he is *S*-able to keep silent; that he is so able is entailed by his keeping silent.¹⁷ And since what we are imagining is a direct “continuation of the present,” it is natural to infer that van Inwagen has that ability already.

Possibly, van Inwagen believes he lacks the ability to keep silent *intentionally*. He may hold that sincerely promising to *A* entails intending to *A*, or entails believing (possibly mistakenly) that one intends to *A*, and he may think that his imagined belief that “the objective, ‘ground-floor’ probability of [his] keeping silent is 0.57” precludes both his intending to keep silent and his believing that he intends to keep silent. Van Inwagen may also think that in the absence of an intention to keep silent, he cannot intentionally keep silent. Alternatively, he may hold that an agent who has only a 0.57 objective probability of keeping silent (given that he promised) lacks sufficient control

over whether he keeps silent to keep silent intentionally. (Compare this agent with someone whose success rate at free throws is 0.57. If, under utterly normal conditions, he sinks his next free throw, is his sinking it an intentional action?)

It also is possible that van Inwagen has in mind a kind of ability that is more reliable than the ability to *A* intentionally. My aim is to locate such a kind of ability. For a time, I use *ensurance-level ability* as its name. In the remainder of this section, I identify and criticize various approaches to locating it. Seeing why these approaches fail will prove instructive.

The control we have over the success of our efforts varies. Michael Jordan has a lot more control over the success of his free throws than Ann does over hers, and Michelle Wie has much more control over the success of her attempts to sink medium-range putts than Bob does over his. Some people may also have more control than others over the success of their efforts to keep silent. One may try to articulate what van Inwagen is after in terms of a high degree of control. It may be suggested that at *t* *S* has ensurance-level ability to bring it about that *p* if and only if it is certain (a "sure thing") that if at *t* *S* were to try to bring it about that *p*, *S* would succeed.

One problem with this suggestion is that cases are imaginable in which although the right-hand side of the biconditional is true, *S* is unable at *t* to bring it about that *p* because he is unable at *t* to *try* to bring *p* about. For example, although it may be certain that if Carl were to try to move his right arm now, he would bring it about that his right arm moves, Carl may be unable to try to move his right arm now owing to hypnosis, and he may now be unable to move it (and to bring it about that it moves) without trying to move it. In such a case, Carl is unable to bring it about that his right arm moves.

Another problem is that we may have ensurance-level ability to bring it about that *p* in cases that have no place for trying to bring it about that *p*. Agents ensure that they intend to *A* in deciding to *A*, since the latter is a mental action of forming an intention to *A*. Possibly, many agents in some ordinary scenarios have ensurance-level ability regarding what they intend. But, as I explained in chapter 2, in normal cases of deciding to *A*, agents do not try to bring it about that they intend to *A*.¹⁸ Nor, in normal cases, does one have an intention to bring it about that one intends to *A* (see chapter 2).

The thesis at issue, again, is this: at *t* *S* has ensurance-level ability to bring it about that *p* if and only if it is certain that if at *t* *S* were to try to bring it about that *p*, *S* would succeed. One might suppose that even though agents who decide to *A* normally do not try to bring it about that they intend to *A*, it is true that if they *were* to try to bring this about they would succeed, and one

of the problems I raised is therefore illusory. This supposition is false. Cases are easily constructed in which someone who decides in the normal way to *A* would not have intended to *A* if he had tried to bring it about that he so intends. For example, in a familiar style of case, the trying would have signaled a manipulator to shut down the agent's brain straightaway.

People talk about doing things "at will." Perhaps ensurance-level ability is associated with that notion. One might suggest that to be able at *t* to *A* at will at *t* is to be so constituted that it is a sure thing that if one were to will at *t* to *A* at *t*, one would *A* at *t*. This suggestion faces a problem of a kind I already mentioned. (If "willing" is another word for trying, it is the *same* problem.) Perhaps Carl is so constituted that it is a sure thing that if he were to will to move his arm now, he would move it now, but, owing to hypnosis, he is temporarily unable to will to move it. Carl may lack even the simple ability to move his arm now, and ensurance-level ability is supposed to be much stronger than simple ability.

One way around the problem is to modify the suggestion as follows: to be able at *t* to *A* at will at *t* is (1) to be so constituted that it is a sure thing that if one were to will at *t* to *A* at *t*, one would *A* at *t* and (2) to be able at *t* to will to *A* at *t*. I must confess that I am unsure what willing is supposed to be. Since I have a firmer grasp on deciding, I rewrite the suggestion in terms of it: to be able at *t* to *A* at will at *t* is (1) to be so constituted that it is a sure thing that if one were to decide at *t* to *A* at *t*, one would *A* at *t* and (2) to be able at *t* to decide to *A* at *t*. A problem with this suggestion is that *deciding* to raise one's arm, say, may itself be something an agent can do at will. Applying the suggestion to a mental action of this kind, we get the following: to be able at *t* to decide at will at *t* to raise one's arm is (1) to be so constituted that it is a sure thing that if one were to decide at *t* to decide at *t* to raise one's arm, one would decide at *t* to raise one's arm and (2) to be able at *t* to decide to decide at *t* to raise one's arm. Needless to say, a commitment to deciding to decide to *A* is to be avoided. (Rewriting the suggestion about willing in terms of *trying* would raise a parallel problem.) Another worry is that libertarians may be entitled to hope for ensurance-level ability even in some indeterministic worlds in which it is never a sure thing that one will act as one decides to act. Perhaps they are entitled to hope for ensurance-level ability with respect to some acts of deciding.

There is a related worry. Suppose that Don's world is deterministic. Don thinks that he has just been given magical powers over dice and he decides to roll a six. Although he lacks special powers, he rolls a six. In a sense, it is a sure thing, prior to his rolling the die, that he will roll a six. That he will roll a six is entailed by a complete description of the state of his universe at any prior

time and of the laws of nature. It also is a sure thing in this sense, prior to his deciding, that he will decide to roll a six then. And, in the same sense, it is a sure thing that if he decides to roll a six, he rolls a six. Depending on how one understands subjunctive conditionals with true antecedents, it may also be a sure thing that if Don were to decide to roll a six (as he does decide), he would roll a six. Even so, Don is not able to roll a six at will. That ability requires special powers, and Don has no such powers.

Perhaps, in ordinary language, the claim that a person is able to *A* "at will" expresses the idea that it is extremely easy for him to *A* intentionally—so easy that his trying unsuccessfully to *A* would be extremely surprising. Perhaps the idea is meant to include the thought that if the person were to try, but fail, to *A*, that would undermine the claim that he was able at the time to *A* at will. It has often seemed to me to be extremely easy for me to *decide* to order beer I like. If I decide to order such beer—a pint of Guinness, say—without trying to decide to order it, then this ease is not properly articulated in terms of trying in the way just identified.

Basic action used to be a hot topic. One might search for ensurance-level ability in that sphere. As I mentioned in chapter 2, a basic action is, roughly, an action that an agent performs, but not by performing another action. My raising my left hand a moment ago was a basic action, if my raising it was an action and I did not raise it by performing some other action—for example, by trying or willing to raise it, where my trying to raise it and my willing to raise it are actions other than my raising it. Again, I am unsure what willing is supposed to be. Whether my trying to raise my hand is an action "other" than my raising it is a subtle question. Perhaps my trying to raise it *is* my raising it, provided that the trying is successful, in which case the fact that I tried to raise my hand does not stand in the way of its being true that my raising my hand was a basic action.¹⁹

Suppose that my raising my left hand was a basic action. Even so, the ability I had to raise it at the time might be less reliable than the ability I had then to perform nonbasic intentional actions of various types. A neurosurgeon might have "randomized" the connection between my acquisitions of intentions (or my tryings) regarding my left hand and bodily motions. Having just acquired the intention to raise my hand (or having just begun to try to raise it), there might have been only a 0.25 chance that things would proceed normally and a 0.75 chance that the result would instead be one of the following: my blinking, my coughing, my sneezing. At the same time, I and my car may be so constituted that my acquiring a proximal intention to start my car would have rendered it virtually certain that I would intentionally start it.

My search for ensurance-level ability thus far has turned up various dead ends. The problems encountered are instructive. Having learned what to avoid, one is in a better position to find what one wants.

4.4. Promise-Level Ability: An Interview

Return to sincere promising. We promise to do things like pay a bill next week and pick up a friend at the airport tomorrow. However, I have never heard anyone promise to *decide* to do such things. A diagnosis of the latter fact is readily available. In promising Bob that I will pick him up at the airport tomorrow, I express an intention to pick him up.²⁰ (If things work out, my intention plays a role in my bringing it about that I do what I promised.) Similarly, a promise now to decide later to *A* would express an intention now to decide later to *A*, and it is only in unusual cases that an agent would intend now to *decide* later to do such things as pick up a friend at the airport or pay a bill.²¹ (He might intend now to decide later *whether* to pick up his friend, but that plainly is another matter.) A philosopher who is guided partly by considerations about the nature of sincere promising, rather than solely by action theoretic considerations, in looking for a kind of ensurance-level ability may be able to set aside general worries about deciding, at least temporarily. The kind of ensurance-level ability I hope to find is associated with sincere promising. I call it *promise-level ensurance ability*, or *promise-level ability* (*P*-ability), for short.

Suppose that Peta sincerely promises to *A*, that she knows what she intends, and that she neither has nor takes herself to have an abnormal source of information about what she will do. (God speaks to van Inwagen in his thought experiment, but he does not speak to Peta, nor does she think he does.) Suppose also that Peta is not up to anything tricky, like intentionally bringing it about that she *A*-s unintentionally (see Mele 1995c, pp. 413–14). Peta may, in fact, unbeknown to her, be unable to *A*. But she does not believe that she is unable to *A*. For given that Peta sincerely promises to *A* and knows what she intends, she intends to *A*, and her so intending is inconsistent with her believing that she is unable to *A*.²² Assuming that Peta is a reflective agent and a fine reasoner, what *does* she believe about her ability to *A*?

A philosopher named Al raised this issue with Peta. He asked first whether she believes that she has nothing more reliable than a *simple* ability to *A*. Peta disavowed that belief, and she pointed out that if she were to believe that, she would lack the confidence that she will *A* that is required for intending to *A*. Al then asked whether she believes that she has nothing more reliable than

the ability to *A intentionally*. Peta disavowed that belief too. On her view, because she is a 90% free-throw shooter, she is able to sink free throws intentionally in normal circumstances. But she does not take herself to be in a position sincerely to promise to sink any of her free throws. Sincerely promising to sink a free throw, Peta said, requires greater confidence that one will sink it than she has, given her knowledge of her success rate.

I return to Al's interview with Peta shortly. A comment on confidence conditions on intending and sincere promising is in order first. Elsewhere, I have defended the thesis that the confidence constraint on intending to *A* is a negative one—roughly, that the agent *not* believe that he will not *A* (Mele 1992a, chap. 8). This constraint will strike some readers as too weak and others as too strong, but there is no need to argue about it here. The point I want to make is that any plausible confidence constraint on sincerely promising to *A* will be stricter: an agent who sincerely promises to *A* believes that he will *A*. In an agent like Peta, that belief is associated with a belief about a very reliable ability, one more reliable than her ability to sink free throws. In this respect, Peta differs from Sue, who also sincerely promises to *A*. Sue believes that God told her that she will *A* if she tries, and she believes, partly on that basis, that she will *A* while also believing that her ability to *A* is limited to simple ability. Here is a concrete illustration. Sue believes that God simply sees that she will sink a free throw straightaway if she tries; he does not, she believes, miraculously beef up her free-throw shooting ability. Sue is fully confident that she will sink her next free throw; and she sincerely promises to do so, even though she knows that her success rate, which she takes to reflect her level of ability, is about 30%.

Return to Al's conversation with Peta. Peta believes that her ability to *A*, which she promised to do, is more reliable than a garden-variety ability to do something intentionally. Al is curious just how reliable she believes it is. He asks whether she believes that the probability of her *A*-ing is 1. Peta replies, "Of course not. As you know, what I promised to do was to meet Pete at the airport early tomorrow morning and drive him home. The airport is ten miles from my house, and I know that things can go wrong on the way. I might be in a serious car accident, for example, or there might be a collision in front of me that blocks the road so long that, by the time I arrive, Pete will have taken a cab home. Other things might go wrong, too. I might need to take one of my kids, or a friend, to the hospital in the morning, my alarm clock might stop working overnight, and so on. I can describe possibilities of mishaps on the way to Pete's house too, but I'm sure you get the point."

Al asks, "So do you believe that you can sincerely make the promise to Pete but not sincerely promise to sink your next free throw because you think your chance of doing what you promised is significantly better than your chance of sinking the shot?" Peta reports that although she does think that she has a better chance of doing what she promised than of sinking her next free throw, subjective probabilities cannot tell the whole story. "Do you know the game Yahtzee?" she asks. "A player throws five dice at a time. Suppose I'd like to roll anything other than five fives—a non-5x5er. My chance of failing to do that, given that I throw the dice, is minuscule. So is the chance of my failing to throw the dice. My chance of failing to roll a non-5x5er is significantly smaller, in my estimation, than my chance of failing to do what I promised Pete I would do. Even so, I am not in a position to promise you that I will roll a non-5x5er."

Al asks why, and Peta replies that he should think in terms of control. She says, "I have no more control over whether I roll a non-5x5er, given that I throw the dice, than I do over whether *you* roll such a roll, given that *you* throw the dice. I cannot literally and sincerely promise anyone that you will roll such a roll, even if I know that you will roll the dice. That is because I have no control over what you roll, given that you roll the dice. (Notice that I potentially do have some control over your rolling the dice. I can offer you a lot of money to roll them.) Together with the comparative point I made, this yields a diagnosis of my not being in a position sincerely to promise to roll a non-5x5er: I can throw the dice, but beyond that I have no control over which spots land face up. To be sure, parents may say such things as 'I promise you that it will rain today' when trying to persuade their children to take an umbrella to school, or, 'I promise you that if you don't drive more carefully, you'll have an accident,' but they aren't speaking literally." "By the way," Peta adds, "I have no more control over which spots land face up, given that I throw the dice, than parents have over the weather."

Because Al suspects that Peta views herself as not being in a position to *intend* to roll a non-5x5er, he sees a potential disanalogy between her Yahtzee scenario, on the one hand, and the free-throw and airport scenarios, on the other. He checks with Peta, who confirms his suspicion. Peta used her Yahtzee example to deflect the suggestion that a difference in subjective probabilities accounts for her belief that whereas she is in a position sincerely to promise to pick Pete up at the airport, she is not in a position sincerely to promise to sink her next free throw. However, Peta lacks an *intention* to roll a non-5x5er, despite her extremely high subjective probability of rolling such a roll. This leaves the following hypothesis open: (*H*) Other things being equal, given

any two courses of action that Peta intends to perform, if she believes that she is in a position sincerely to promise one but not the other, that is because of a significant difference in subjective probability of success. This hypothesis is associated with a simple idea about the difference between *P*-ability and *I*-ability: (*Simp*) Regarding intended actions of kinds the agent often performs, what separates *P*-ability from *I*-ability is simply a significant difference in relevant success rates.

A straightforward test of hypothesis *H* compares relevant cases in which Peta's subjective probability of success regarding intended courses of action is the same. Imagine now that Peta is an extraordinarily accomplished, 98% free-throw shooter, as she knows. She also has, as she knows, a 98% success rate at fetching people she intends to fetch from the local airport. Such fetching is part of her job, and she has done this hundreds of times in the past several years.

Some readers may worry that promising is inappropriate in the free-throw case for a special reason linked to the point that, normally, the only permissible goal of a player at the free-throw line is sinking the shot.²³ I circumvent this worry by supposing that Peta is playing a game in which one announces one's goal at the line. Permitted goals include sinking the shot and missing it by deflecting the ball off of an announced part of the rim (left, right, front, or back). Peta has played this game a lot and has, as she knows, her normal 98% success rate in it of sinking intended free throws.

Can Peta reasonably and correctly believe that although she is in a position sincerely to promise to pick Pete up at the airport, she is not in a position sincerely to promise to sink the free throw that she intends to sink now? Those who judge that the answer is yes probably will find *Simp* too simple. If Peta correctly believes the proposition at issue, a plausible diagnosis of the correctness of her belief includes the judgment that it is false that what separates *P*-ability from mere *I*-ability (in cases of the sort at issue) is a disparity in relevant success rates. What about those who judge that the answer is no? They may find *Simp* attractive. Each group may draw a distinction between *I*-ability and *P*-ability. But how should my question about Peta be answered?

My aim is to distinguish *P*-ability from *I*-ability in a way that is sensitive to commonsense judgments. Now, people routinely sincerely promise to fetch others from airports, and although they do not assign precise probabilities to their being successful, a subjective probability of 0.98 would seem not to be far off a normal person's actual mark and would seem not to preclude sincere promising.²⁴ (Of course, I have in mind only people who have done a lot of airport pickups.) If there were 98% free-throw shooters with a good grip

on the concept of promising who played games like the one I made up and were sometimes asked (by teammates, for example) for promises to sink shots, would they sometimes respond with sincere promises to sink their shots? It is hard to say. Peta reports that even if her success rate were 98%, she would feel too uncomfortable about making such promises to make them, owing to her imperfect control over relevant bodily events that partly constitute her free throws. Typically, she says, her misses feel just like her successful shots. Try as she may, she says, she cannot shrink her 2% margin of internal error. Because this margin of error remains, Peta reports, and because its source is internal, she would feel extremely uneasy about promising. Peta says that when an equally small margin of error derives from such external factors as unexpected traffic conditions or car failures, she has no qualms about promising. I return to this issue in section 4.5.

Peta has done enough work. She lacks the patience for various further subtleties. One might suggest that if people were to realize that they rarely can be fully confident that they will do the things they promise to do, they would make very few of the promises they do, and that sincere promising requires greater confidence than Peta has that she will fetch Pete from the airport. An alternative suggestion is that when people say such things as “I promise to meet you at the airport,” what they really mean is that they promise to make a genuine effort to do that, unless they acquire a very good reason for not meeting the person or become incapable of meeting him. One who makes the latter, deflationary suggestion may also claim that people are entitled to be extremely confident that they will keep such promises, and significantly more confident than Peta is about picking up Pete. Perhaps close attention to promising would provide significant support for one of these suggestions, and perhaps not.²⁵ It suffices for immediate purposes to notice that there is a clear difference between *S*-ability and *I*-ability and a *prima facie* difference between both of these abilities and an ability that sincere promise-makers like Peta have, if things are as they take them to be. This is consistent with the suppositions that these abilities lie on a continuum, that the boundaries are fuzzy, that there are intermediate abilities, and that there are stronger abilities than *P*-ability.

4.5. Promise-Level Ability Pursued

One way to approach promise-level ensurance ability (*P*-ability) is to ask what agents with a firm grip on the concept of promising who have no abnormal source of beliefs about what they will do must believe or

presuppose about their ability to *A* in order sincerely to promise to *A*. That is why I staged the interview with Peta. My reason for concentrating on agents who understand promising is obvious. Agents who are confused about promising are not reliable guides for present purposes. I concentrated on an agent with no abnormal source of beliefs about what she will do because some beliefs with abnormal sources, as I have explained, may support a sincere promise even in agents who realize they have only minimal relevant abilities. Again, I am after a kind of ability to *A* that is stronger than the ability to *A* intentionally.

I restrict the ensuing discussion in a related way. Elsewhere, I have argued that an agent can intentionally bring it about that he *A*-s unintentionally (Mele 1995c, pp. 413–14), and I have defended the possibility of an agent who intends to *A* while knowing that if he does not *A* intentionally, he will *A* unintentionally (Mele 1992b). I want to avoid complexities that such cases would introduce into the present discussion. This is not because the complexities are worrisome, but rather because dealing with them would require considerable space and take the spotlight off the central point of interest—a practical ability, associated with promising, that is stronger than *I*-ability. Now, promisers do not say such things as “I promise to pick you up at the airport intentionally.” Nor do they promise to pick us up unintentionally, or nonintentionally, or perhaps intentionally but perhaps not. Ordinary promise makers make their promises and leave the action-theoretic work to us. In stereotypical cases of sincere promising, the promised prospective course of action—the course of action the agent represents herself to the promisee as intending to take—is an intentional one. In pursuing *P*-ability I consider only agents who promise to *A* and disbelieve all of the following: that they will *A* unintentionally; that they will *A* nonintentionally; that they will *A* but perhaps intentionally and perhaps not.²⁶

Another restriction on the agents I consider is motivated partly by the point that people who discover that they have been tricked into making a promise are often within their rights to refuse to keep it.²⁷ They often have a good excuse for not keeping such a promise. The further restriction is that the agents persist in believing that they were not manipulated into making the promise and in believing that they did not mistakenly make it. In sum, I investigate *P*-ability by considering agents with a firm grip on the concept of promising and no abnormal source of beliefs about what they will do who promise to *A*, persist in believing that they were not manipulated into so promising and that they did not mistakenly so promise, and disbelieve all of the following: that they will *A* unintentionally, that they will *A* nonintentionally,

that they will *A* but perhaps intentionally and perhaps not. I call such agents “*C* agents.”

Here is a simple hypothesis. In order to make a sincere promise to *A*, *C* agents must believe or presuppose the following: (*Ar*) It is extremely likely that if they promise to *A*, they will *A*. As I understand it, this belief or presupposition condition is meant to be stronger than that for intending to *A*. For example, Peta, a 90% free-throw shooter, may intend to sink her next free throw without believing or presupposing that it is *extremely* likely that she will sink it or that it is *extremely* likely that she will sink it if she intends to sink it. This simple hypothesis coheres with *Simp* (in section 4.4).

As I have mentioned, Peta has reservations about these ideas. Readers who share them will be dissatisfied with the simple hypothesis and *Simp*. Here is a hypothesis for such readers. In order to make a sincere promise to *A*, *C* agents must believe or presuppose something to the following effect: (*Br*) Their ability to *A* is such that they are entitled to be fully confident that, barring unexpected substantial obstacles, if they sincerely promise to *A*, they will *A*.²⁸ This would explain why Peta does not take herself to be in a position sincerely to promise to sink her next free throw (given either her actual success rate or the imagined 98% success rate). Peta is not—nor is she entitled to be—fully confident that, barring unexpected substantial obstacles, if she sincerely promises to sink the free throw, she will sink it. She knows that her control over the success of her attempts—her general free-throw shooting ability—does not warrant full confidence in this. However, she believes that her relevant abilities are such that she is entitled to be fully confident that, barring unexpected substantial obstacles, if she sincerely promises to fetch Pete from the airport, she will do so.

Obviously, I am assuming that, normally, when Peta misses a free throw, her failure is not due to her encountering an unexpected substantial obstacle. Substantial obstacles include such things as sudden cramps or vertigo, blurred vision, and a fan’s shooting the ball with an arrow in mid-flight. They do not include small-scale bodily events that are in Peta’s normal range when shooting free throws but sometimes add up to her releasing the ball a little too early or too late, pushing a bit too much or too little with her legs, or the like (see section 4.2). Even extraordinary free-throw shooters are not as reliable at sinking their free throws *in the absence of unexpected substantial obstacles* as many ordinary folks are entitled to count on themselves to be at picking up friends at airports *in the absence of such obstacles*.

As I understand unexpected substantial obstacles, they are unexpected *by the agent* and an agent cannot expect to encounter unexpected obstacles. To be

sure, assertions like the following are intelligible: "I always encounter bizarre obstacles I don't expect when I sail through the Bermuda Triangle, so I expect to encounter unexpected obstacles—specific obstacles I don't expect—this time too." However, in *Br* and subsequent discussion, the expression is used generically. The idea, more clumsily expressed, is that in order to make a sincere promise to *A*, *C* agents must believe or presuppose something to the following effect: (*Br**) Their ability to *A* is such that they are entitled to be fully confident that, if, as they expect, no substantial obstacles to their *A*-ing arise (or exist already), they will *A* if they sincerely promise to *A*.

Here is another hypothesis (and an apparent truth) about promising. In order to make a sincere promise to *A*, *C* agents must believe or presuppose something to the following effect: (*B*₂) Barring unexpected substantial obstacles that they would reasonably take to warrant abandoning their intention to *A*, if they sincerely promise to *A*, they will not abandon their intention to *A*.²⁹ On the following grounds, I take *B*₂ to be implicit in *Br*. Can a *C* agent who is *doubtful* about a first-person instance of *B*₂ consistently believe a related first-person instance of *Br*? Not as I understand *Br*. As I understand *Br*, (1) that *one's ability to A is such* that one is entitled to be fully confident that, barring unexpected substantial obstacles, if one sincerely promises to *A*, one will *A* entails (2) that one is entitled to be fully confident that, barring such obstacles, if one sincerely promises to *A*, one will *A*. And a *C* agent who is doubtful about a pertinent first-person instance of *B*₂ cannot consistently believe 2. Thus, I understand Peta's ability to fetch Pete from the airport, for example, to encompass an ability to resist temptations to abandon, against—or without the support of—her better judgment, an intention to do that.³⁰

In an effort to locate promise-level ability to *A*, I have been discussing something that *C* agents must *believe* or *presuppose* about their abilities in order to promise sincerely to *A*. Here is a related hypothesis about promise-level ability itself:

P. *X* is a promise-level ability to *A* only if *X* is a sufficiently reliable ability to ground, in a *C* agent who knows his own abilities, complete confidence that, barring unexpected substantial obstacles, if he sincerely promises to *A*, he will *A*.³¹

This, of course, is a statement of an alleged necessary condition for something's being a *P*-ability to *A*, not a statement of necessary and sufficient conditions. My concern was to find a kind of ability more reliable than a garden-variety ability to *A* intentionally. *P* identifies a mark of such a kind of ability. Often,

agents who *A*-ed intentionally, and therefore were able at the time to *A* intentionally, lacked an ability to *A* with the kind of reliability mentioned in *P*. Just think of all those very good free-throw shooters, golfers, and eight-ball players who intentionally sink relatively easy shots that they are in no position sincerely to promise to sink. Their pertinent abilities are not sufficiently reliable to ground, in a *C* agent who knows his own abilities, the confidence specified in *P*.

Although I do not try to augment *P* to generate a statement of necessary and sufficient conditions for something's being a *P*-ability, an issue central to that project should be identified. Suppose a *C* agent, Cam, believes that there is about a 20% chance that something unexpected will prevent her from picking up Bob at the airport tomorrow morning. Cam may be completely confident that, *barring unexpected substantial obstacles*, if she sincerely promises to pick Bob up, she will do so. But can she sincerely promise to pick him up? The intuitive answer is *no*. Seemingly, sincerely promising to *A* requires that one not believe that one's chance of *A*-ing, even if one does one's best to *A*, is only about 0.8. Must Cam believe that there is *no* chance that something will prevent her from picking Bob up in order to be in a position sincerely to promise to pick him up? Not if normal agents—who realize that there is some chance of failure—are often in a position sincerely to promise to do such things as pay their bills and fetch others from airports. These points about beliefs suggest that an augmented version of *P* will include a clause requiring that the chance of unexpected substantial obstacles not be too great without requiring that it be 0. I do not speculate further about such a clause.

A comment on my strategy in sections 4.4 and 4.5 is in order. It is a datum that people do not take themselves to be in a position sincerely to promise to do some ordinary things that they take themselves to be in a position to intend to do and to be able to do intentionally. This datum, *D*, may be interpreted in light of two others. At least in normal scenarios, the kind most relevant to the present inquiry, (1) anyone in a position sincerely to promise to *A* is in a position to intend to *A*, and (2) anyone who takes himself to be in a position to intend to *A* takes himself to be able to *A* intentionally. A plausible hypothesis about *D*, in light of the other data, is that typical promise-makers have, at least tacitly, the view that sincere promising (or perhaps paradigmatic sincere promising) requires a higher estimation of one's abilities than intending does, or at least make sincere promises in a way that coheres with this view. My primary concern is the relevant abilities themselves, not agents' beliefs or presuppositions about their abilities, not fine points about promising, and

not the ability to make promises. I have been attending to (paradigmatic) sincere promising to help me locate a kind or level of ability that is more robust than *I*-ability. For the purposes of this chapter, my interest in promising is purely instrumental.

Both *P* and *Simp* (the hypothesis that in the case of intended actions of kinds the agent often performs, what separates *P*-ability from *I*-ability is simply a significant difference in relevant success rates) are meant to be sensitive to intuitions about cases. Does commonsense say, with Peta, that in scenarios in which a player's extraordinarily high free-throw percentage matches her success rate at fetching people from airports (which she does frequently), she is in a position sincerely to promise to fetch a friend from the airport but not in a position sincerely to promise to sink her next free throw, given that she knows her success rates and understands what promising is? Does it say instead, with *Simp*, that under these conditions and other things being equal, if she is in a position sincerely to promise either, she is in a position sincerely to promise the other? If the difference between having a 2% failure rate owing to unexpected external obstacles and having the same failure rate owing to imperfect control over internal, bodily events that partly constitute one's free throws makes no difference in the present context, *P* should be abandoned. If, as Peta thinks, the difference does make a difference, *Simp* is false.

I contend that *P* is more strongly supported than *Simp* by a commonsense view of sincere promising. *P*, unlike *Simp*, is suggestive of an element of self-trust that (paradigmatic) sincere promising encompasses. In paradigmatic cases at least, sincere promisers trust themselves to do what it takes to keep their promises if unexpected excusing conditions do not arise. There is no hint of this in *Simp*. Now, being an imperfect free-throw shooter is no excuse for missing a free throw, as any fan will tell you. Excuses include items on my partial list of substantial obstacles—for example, cramps and blurred vision—but not relatively tiny muscular events that are in the agent's normal range when shooting free throws, even though those events occasionally combine to yield a failed attempt. Even the imaginary Peta who is a 98% free-throw shooter under normal conditions misses 2% of her shots under such conditions. When she fails under these conditions no excuse is available to her. Perhaps it is because Peta realizes this that she does not take even her imaginary extraordinary self to be in a position sincerely to promise to sink a free throw. The connection between promising and excuse supports *P* and not *Simp*.³²

4.6. Recap and an Attempt at Provocation

I have explored three kinds or levels of practical ability: simple ability to A (S -ability), ability to A intentionally (I -ability), and what I called promise-level ability to A (P -ability). I have pointed out that the second kind of ability is stronger than the first, in the sense that it entails, but is not entailed by, the first. Any agent who is able to A intentionally is able to A ; but an agent who is able to A may not be able to A intentionally. In the case of action-types that are *essentially* intentional, like *deciding* to order a Guinness or *trying* to bend one's elbow, to be S -able to perform an action of that type is to be I -able to do so.

Given normal and even some stricter standards for intentional action, we are able to do many things intentionally that we lack promise-level ability to do: recall those very good golfers and free-throw shooters who intentionally sink putts and shots that they are in no position sincerely to promise to sink. So I -ability does not entail P -ability. And there is at least a sphere in which P -ability is sufficient for I -ability. Having promise-level ability entails being able to do what one promised. All sincere promises to A made by C agents are, at least tacitly, promises to A intentionally. So, in C agents who promise to A , being P -able to A suffices for being I -able to A .

My chief hope for this chapter is that it will motivate at least a few of us to keep in mind that there are the different kinds or levels of practical ability examined here when we explore such questions as whether free will and moral responsibility require the ability to do otherwise, what we are able to do if we have the power of agent causation or various indeterministic powers, whether "ought" implies "can," what it is for a desire to be irresistible, and whether a person who akratically experiments with crack cocaine is able to exercise self-control in a way that some crack addicts are not. I believe that our bearing in mind differences among these three kinds of ability will improve our discussions of such topics. For example, if van Inwagen had attended to these different kinds of ability, he would not have claimed that he was not able to keep silent, although he might have claimed that he was not able to keep silent intentionally or, instead, that he had the latter ability but lacked a promise-level ability to keep silent. That no agent who knows that it is undetermined whether he will keep silent about a friend is S -able to keep silent is an astounding result, one that would indeed make free will, on libertarian views, look mysterious, as van Inwagen claims it is. But van Inwagen does not produce this result. Nor does he produce the result that no agent who knows this about himself is I -able to keep silent. Just imagine an agent who does

know this about himself and whose chance of keeping silent if he intends to do so is as high as Peta's chance (under normal conditions) of sinking her next free throw, as she intends to do; and imagine that his confidence about this matches Peta's about her shot and that he intends to keep silent. Furthermore, even if we were to be persuaded that this agent—and any agent who knows that it is undetermined whether he will keep silent even if he intends to—lacks *P*-ability to keep silent, we would need to think hard about what implications this lack of *P*-ability would have for libertarianism.

Some readers may claim (uncharitably!) that although attention to different levels or kinds of practical ability would have helped van Inwagen, it would not be generally useful in exploring any of the questions I identified. I conclude with a brief reaction to that claim, in the form of an illustration. Here is a conjecture: if an agent's freely *A*-ing at *t* requires his being able at *t* to perform an action that is an alternative to *A*, the level of the required "alternative" ability is no higher than the highest-level ability to *A* required for his freely *A*-ing. The level at which that ability lies may vary depending on the kind of action at issue (for example, on whether it is a basic or nonbasic overt action or a decision). Consider a common kind of action—voting. In order to have voted freely for Gore, must Al have been either *P*-able or *I*-able to vote for him? Well, here are the facts about Al (see Mele 1995a, p. 14, n. 11). Intending to vote for Gore, he pulled the Gore lever in a Florida voting booth. Unbeknown to Al, that lever was attached to an indeterministic randomizing device: pulling it gave him only a 0.001 chance of actually voting for Gore. Luckily, he succeeded in registering a Gore vote. Beyond the rigging of the voting booths at Al's voting establishment, there is no monkey business in Al's story. He is not brainwashed, for example. And Al is a sane, rational adult whose intention is backed by reasons he had for voting for Gore. Moreover, by hypothesis, free actions are common in Al's world.

It is very plausible that Al's voting for Gore (which, as I understand it, requires actually registering a Gore vote) was too lucky to count as an intentional action (see Mele and Moser 1994) and that, given his circumstances, Al was not *I*-able (hence, not *P*-able) to vote for Gore at the time. However, if free actions are common in Al's world, it is difficult to see why his voting for Gore should not count as a free action, other things being equal. If the action is free and if what I said is very plausible is true, Al freely voted for Gore while being neither *P*-able nor *I*-able to vote for him. And, in that case, if my conjecture also is true, any ability to perform an alternative action to voting for Gore that Al might have needed to vote freely for Gore is weaker than *I*-ability. So if, as some theorists hold, Al's having freely voted for Gore at *t* requires that

he was able at t to do otherwise than vote for Gore, how is that ability to be understood? Is S -ability enough? Does Al need something stronger than that but weaker than I -ability? It certainly looks like attention to levels of practical ability is in order.

It may be replied that my conjecture is what generates this appearance and the conjecture is false.³³ Readers who find that reply attractive are invited to argue for it without attending to different levels of practical ability. And all readers are encouraged to reflect on whether van Inwagen was right to emphasize, in his criticism of agent causationists, whatever kind of ability he had in mind, given that Al freely voted for Gore even though he was neither P -able nor I -able at the time to vote for him. Perhaps van Inwagen was wrong to emphasize what he did, and perhaps not. That depends partly on whether there are actions that differ from Al's voting for Gore in such a way that a significantly higher level or more robust kind of ability is required for freely performing them.³⁴

Notes

1. Although I am not able to golf just now, or to golf two minutes from now, I am able to get to a driving range in about twenty minutes. It is very natural to say that I am able now to start hitting golf balls in twenty minutes or so.
2. J. L. Austin writes, "of course it follows merely from the premise that he does it, that he has the ability to do it, according to ordinary English" (1970, p. 227).
3. Tomis Kapitan notes a similar distinction between abilities (1996, pp. 102–4).
4. See Mele 1995a, pp. 211–21. Also see John Fischer's distinction between "guidance" and "regulative" control (1994, pp. 132–35).
5. In ordinary English, people sometimes balk at moving from "able" claims to corresponding "ability" claims. Ann rolled a six with a fair die. It is natural to say that she was able to do that, and it is perhaps less natural to say that she had an ability to do that. However, notice the awkwardness of the following assertion: "Ann was able to roll a six, but she had no ability to roll a six." Of course, a speaker who makes this assertion can draw a distinction in light of which what he means to assert is true. For example, he can say that he understands " S was able to A " in such a way that it is entailed by " S A -ed" and that he understands having an ability to A as entailing being able to A intentionally. Given this chapter's purpose, attention to alleged differences between "able" claims and "ability" claims would be a source of distraction.
6. Libertarians who reject the idea that there are indirectly free actions should ignore the word "directly" in "directly free."
7. A comment on time t is in order. Some actions take more time than others to perform. In the case of a nonmomentary action A performed at t in W , the

possible worlds at issue have the same laws of nature as *W* and they have the same past as *W* up to a moment at which the agent's conduct first diverges from his *A*-ing. This initial divergence can happen at a moment at which the agent is *A*-ing in *W* or at the moment at which his *A*-ing begins in *W* (see Mele 2006, pp. 15–16).

8. The "*L*" stands for "libertarian," since libertarians and other incompatibilists typically favor an understanding of ability along these lines. The analysis offered of simple *L*-ability can be strengthened as follows for a libertarian who holds that even an agent who *A*-ed (intentionally) was not able to *A* unless he was also able at the time not to *A*: *S* has, at the relevant time, the simple *L**-ability to *A* at *t* if and only if either (1) *S A*-s at *t* and there is a possible world with the same past and laws in which *S* does not *A* at *t* or (2) *S* does not *A* at *t* and there is a possible world with the same past and laws in which *S A*-s at *t*.
9. Freedom-level ability may be understood as a kind of ability such that if, setting aside ability conditions, everything necessary for an action's being free were present, adding a suitably exercised ability of this kind would yield sufficient conditions for the action's being free.
10. "*B*," like "*A*," is to be read as an action variable. Again, I do not take a stand on how actions are to be individuated—for example, on whether Fred's rolling the die and his rolling a five are the same action under different descriptions or different actions. (If Fred's rolling the die and his rolling a five are the same action under different descriptions, the same action can be intentional under one description and not intentional under another.)
11. For stylistic reasons, I will stop mentioning counterparts. Readers who reject the idea that the same agent can be located in different possible worlds should henceforth make the relevant substitutions. For example, in stories in this book in which the same agent is located in pairs of possible worlds, they should regard the agents as counterparts.
12. The difficulty of producing an *analysis* of ability from this perspective has been a thorn in the side of traditional compatibilists, who agree with libertarians that freely *A*-ing and being morally responsible for *A*-ing require that one is able to do otherwise than *A* but disagree about the nature of this ability.
13. Semicompatibilism is sometimes misrepresented as compatibilism about moral responsibility and incompatibilism about free action, an issue I discuss in chapter 5, section 5.3.
14. Whether one would say that Shaq is *I*-able to sink his free throw depends on one's view about whether agents are *I*-able to do things they succeed at doing on about half of their attempts.
15. Agent causation is characterized and discussed in subsequent chapters.
16. Van Inwagen's claim about agent causation is that the further knowledge that he "will be the agent-cause" of his conduct in this scenario would not undermine his belief that he is not able to keep silent (2000, p. 18).

17. The abilities that concern me in this chapter, as I said, are *actional* ones. It is not clear that keeping silent is an action, even when one intentionally keeps silent (see Mele 2003a, pp. 146–64). For the purposes of this chapter, however, the simplifying (and, in my view, false) assumption that all intentional “not-doings” (e.g., not telling on one’s friend, not voting in today’s election) are actions is harmless.
18. This is not to say that deciding is effortless. Perhaps, in deciding to *A*, one normally is trying to settle some practical question or other—for example, “What shall I do?” or “Shall I *A* or *B*?”
19. See Adams and Mele 1992, pp. 329–30. Also see Hornsby 1980.
20. As I use “express an intention,” one may express an intention that one mistakenly believes one has.
21. For an unusual science fiction case of this kind, see chapter 2, section 2.2.
22. It may be claimed that strange agents sometimes believe that they are able to *A* while also believing (without equivocation) that they are unable to *A*. Peta is not strange in this way.
23. Basketball fans know that in special situations other goals make strategic sense—for example, deflecting the ball to a teammate who can take a three-point shot.
24. Athletes occasionally “guarantee” that their teams will win their next game. This sounds a bit like promising, but thoughtful auditors realize that the players are not speaking literally and that they would not be speaking literally if they were to say “promise” rather than “guarantee.”
25. “Perhaps not” is too modest in my opinion. But there is no need to defend that opinion here.
26. On a middle ground between intentional and unintentional action—nonintentional action—see Mele 2012b.
27. I am grateful to Pekka Väyrynen for this observation years ago.
28. As I have implied, not all obstacles are brutally physical. As I use “obstacle,” that Peta’s child needs to be taken to the hospital in the morning is an obstacle to her picking up Pete at the airport, as promised.
29. Here is a formulation of *B*₂ to match *B*₁*: (*B*₂*) Barring substantial obstacles at odds with their expectations that they would reasonably take to warrant abandoning their intention to *A*, they will not abandon their intention to *A* if they sincerely promise to *A*.
30. When Peta promises to pick Pete up at the airport, she generates a reason to pick him up. Prior to promising, she may have had good reasons to pick him up, reasons having to do with their friendship. But, having promised, she has even better reasons to do so. The extra reason created by her promise might give Peta a higher threshold for intention-abandonment than she would have had if she had intended, but not promised, to pick Pete up. It may be that some unexpected occurrences that she would regard as warranting abandoning her intention in the latter scenario, she would not so regard in the actual scenario, given her promise. Thus, her promise may make it more likely that Peta will pick Pete up than would have been the case if she had intended, but not promised, to do so. However, this

is not a consideration I wish to highlight; for the difference between *I*-ability and *P*-ability does not lie here. No matter how firmly a 90% free-throw shooter intends to sink her next free throw, she is much more likely to miss the shot in the absence of unexpected substantial obstacles than an ordinary person who equally firmly intends to pick up a friend at an airport is to fail to pick the friend up in the absence of such obstacles. That is because there is a significant difference in control in the two cases, a difference that helps explain why it is that we are entitled to be fully confident that we will succeed in driving across town, *barring unexpected substantial obstacles*, but even great free-throw shooters are not entitled to be fully confident that they will sink their next free throw, *also* barring such obstacles. (I am grateful to Gideon Yaffe for encouraging me to make this point years ago.)

31. Obviously, *P* does not entail that only *C* agents have *P*-abilities. Abel is not a *C* agent, but he has an ability to *A* that is sufficiently reliable to ground in a possible *C* agent who has Abel's ability to *A* and who knows his own abilities complete confidence of the kind specified in *P*. Here is a formulation of *P* to match *Br*: (*P**) *X* is a promise-level ability to *A* only if *X* is a sufficiently reliable ability to ground, in a *C* agent who knows his own abilities, complete confidence that, if, as he expects, no substantial obstacles to his *A*-ing arise (or exist already), he will *A* if he sincerely promises to *A*. The point just made also applies to *P*.*
32. Imagine a 98% free-throw shooter, Lita, who knows that she misses free throws when and only when a certain twitch occurs in her right wrist during a shot. Because the twitch occurs in only 2% of her attempts, she always expects it not to occur. Can Lita, who understands promising, sincerely promise to sink her next free throw? Perhaps, owing to the remoteness of this scenario from ordinary free-throw shooting, commonsense yields neither a yes nor a no answer. To the extent to which one sees the twitch as similar to unexpected external events that would provide effective excuses for not doing what one promised, one may view Lita as being in a position sincerely to promise to sink her next free throw. If Lita cannot sincerely promise to sink it, a statement of necessary and sufficient conditions for *P*-ability should include a condition that entails that Lita's ability to sink her next free throw is not a *P*-ability. Again, my aim in this chapter does not include providing an *analysis* of *P*-ability.
33. My conjecture is just that. I am not claiming that it is true.
34. Of course, deciding to *A* requires a higher level of ability, since such deciding is essentially intentional. My point is that a proper investigation of the question whether van Inwagen was right or wrong to emphasize what he did will be sensitive to levels or kinds of practical ability. For comments on a draft of Mele 2003b, on which this chapter is based, I am grateful to Helen Beebe, Randy Clarke, Josh Gert, Alan Goldman, Risto Hilpinen, Jamie Hobbs, Cei Maslen, Michael McKenna, Eddy Nahmias, Dave Robb, Pekka Väyrynen, David Widerker, Gideon Yaffe, Aaron Zimmerman, and audiences at Cornell University, Florida State University, the University of Colorado at Boulder, and the University of Miami.

Free Will and Moral Responsibility

DOES EITHER REQUIRE THE OTHER?

HOW ARE FREE will and moral responsibility related to one another? In this chapter, I investigate four more-specific questions on this topic in the sphere of action:

1. Can agents be morally responsible for *A*-ing even though they did not freely *A*?
2. Can agents freely *A* without being morally responsible for *A*-ing?
3. Can agents who are never morally responsible for anything sometimes act freely?
4. Can agents who never act freely be morally responsible for some of their actions?

5.1. Warming Up

Does Al deserve some credit or blame from a moral point of view for voting for Gore in the story I told about him toward the end of the preceding chapter? If so, he is morally responsible for voting for Gore, as moral responsibility is often understood. Assume that moral responsibility is common in Al's world. Then, in my view, he deserves some moral credit for voting for Gore and so is morally responsible for doing that. He is morally responsible for the action in what is now sometimes called the “accountability” sense (Shoemaker 2011; Watson 1996)—the only sense of “moral responsibility” that directly concerns me in this book.

Did Al freely vote for Gore? Was his voting for Gore a free action? Did he vote for Gore of his own free will? These are three ways of asking the same question (at least as I am using the relevant terms). Assume that free actions are common in Al's world. Even then, someone might claim, he does not freely vote for Gore. Someone who has heard it said that freedom is the control condition on moral responsibility and finds that plausible may think that Al does not have enough control over whether he votes for Gore to vote for him freely. If such a person shares my opinion that Al is morally responsible for voting for Gore, he or she may regard my story as an example of an agent who is morally responsible for doing something that he does not freely do.

My friend Ann, who has never studied philosophy, says that she has no doubt that Al freely votes for Gore. Keeping in mind the assumption that free actions are common in Al's world, what may be said in support of Ann's opinion? If we understand trying to *A* in an unexacting way that is popular in the philosophy of action literature, we should say that Al tried to vote for Gore. As I observed in chapter 3, trying to *A*, on the conception of it at issue, requires no *special* effort. For example, when I turned my computer on this morning, I tried to turn it on, even though I turned it on simply by pressing a button. I expended very little energy and very little effort, but trying to turn on my computer does not require much of either. Now, if free actions are common in Al's world, then it is plausible that he freely *tried* to vote for Gore, given the details of the case. And Al's voting for Gore may count as a free action in virtue of its relationship to his free attempt to vote for Gore. If his so doing is properly counted as a free action on these grounds, then his voting for Gore may be said to inherit its status as a free action from a free action in which his trying to vote for Gore partly consists—that is, from his pulling the “Gore” lever, something he did with the intention of voting for Gore. We can say that his pulling the Gore lever is a directly free action and that his voting for Gore is indirectly or derivatively free.

Deciding to vote for Gore is no part of my story about Al. Decisions to do things, as I understand them, are responses to uncertainty about what to do (chapter 2). If Al was at no point uncertain about whether to vote for Gore (if voting for Gore was, as we say, a no-brainer for him all along), then there is no place in Al's story for his making a decision to vote for Gore. But imagine another voter, Betty, who is uncertain about whom to vote for, and even about whether to vote at all, and eventually decides to vote for Gore. Betty votes at the same place as Al, has the same chance of actually producing a Gore vote by pulling the Gore lever, and, like Al, luckily succeeds in voting for

Gore. Also, like Al, she is rational, unmanipulated, and so on. If Betty freely votes for Gore, her voting for him may be said to inherit its status as a free action at least partly from her freely *deciding* to vote for him. Her pulling the Gore lever may also be said to inherit its status as free from her freely deciding to vote for him. We can say that her deciding to vote for Gore is a directly free action and that her pulling the Gore lever and her actually voting for him are indirectly free.

5.2. The Easier Questions

In the preceding section, I entertained the idea that an agent may be morally responsible for an action of his that he did not freely perform. I do not regard my stories about the Florida voters as cases in point. But other stories might work. I continue to assume for a while that free actions and actions for which agents are morally responsible are common in relevant possible worlds. Consider the following story. Van, a normal man, got drunk at a party and then tried to drive home. He was so drunk that he did not realize he was impaired. No one tricked Van into drinking alcohol, no one forced him to drink, he is knowledgeable about the effects of alcohol, and so on. Owing to his drunkenness, he accidentally drove into and killed a pedestrian he did not see.

It is commonly claimed that, in cases of this kind, the driver is morally responsible for killing the pedestrian. And I agree. However, the claim that Van freely killed the pedestrian is difficult to take seriously. Van was not trying to kill the pedestrian, he was not trying to harm anyone at the time, he had no idea a pedestrian was in his path, he killed the pedestrian accidentally, and so on.

This example speaks in favor of a yes answer to my question whether an agent who does not freely *A* can be morally responsible for *A*-ing. Does it have a bearing on my question whether agents who *never* act freely may be morally responsible for some of their actions? A common response to stories like Van's is that the driver is indirectly morally responsible for killing the pedestrian partly because he is directly morally responsible for some relevant *free* action (for example, drinking excessively). This response represents Van's moral responsibility for the killing as dependent on an earlier free action. If Van had been force-fed alcohol and placed in the driver's seat of his car and was so drunk that he did not realize he was impaired, we would not see him as morally responsible for the killing. According to the common view I mentioned, that is because there is no relevant free action by Van in the chain of events

leading up to the killing. My question about agents who never act freely needs separate treatment. I take up that question in sections 5.3 and 5.4.

The next question on the agenda is whether agents can freely *A* without being morally responsible for *A*-ing. As I understand moral responsibility, it is a moral matter. So, in my view, agents are not morally responsible for actions that fall outside the sphere of morality—or, more precisely, actions that morality is not in the business of prohibiting, requiring, or encouraging.¹ Consider the following case. Carl, who is taking an undergraduate course on free will, has been thinking about free will a lot lately. There is a tree in the middle of the sidewalk he takes from his physics class to his philosophy class, and there is a circular path around the tree. As Carl approaches the tree, he thinks about spontaneously selecting one of the two normal routes past it. He wonders what a conscious spontaneous decision about this would feel like. He decides to pass the tree on the left. Carl is sane, rational, reasons-responsive, unmanipulated, and so on.

Although I am no expert on morality, I do not feel at all sheepish about counting Carl's decision to go left as falling outside the sphere of morality (in the sense identified above) and therefore as one to which moral responsibility does not apply. Even if I am right about this, is this action a free action (assuming that free actions are common in Carl's world)? As some philosophers conceive of free will (Campbell 1957, pp. 167–74; Kane 1989, p. 252), exercises of it can occur only in situations in which people make important moral or practical decisions in the face of temptation or competing motivation.² This obviously is a far cry from what Carl did; he attaches no special importance to either path. However, other philosophers are much less restrictive about free will (Clarke 2003, chap. 7; Fischer and Ravizza 1992; O'Connor 2000, pp. 101–7). For readers who believe that Carl's decision is a free action, a yes answer to the present question should sound right—provided that they agree with me that Carl's decision is outside the sphere of morality.

The third question about the relationship between free action and moral responsibility that I raised is whether it can happen that agents who are never morally responsible for anything sometimes act freely. If the answer to this question is yes, then, of course, the answer to the question whether it can happen that agents are not morally responsible for *some* of their free actions is yes as well.

In an article on religion and morality, George Mavrodes suggests that morality is “provisional and transitory, that it is due to serve its use and then to pass away” (1986, p. 226). Perhaps there is a respectable conception of some intelligent heavenly agents according to which morality no longer applies to

their actions. If there can be agents who are at least as intelligent as the average adult human being but to whose actions morality does not apply, might they sometimes act freely?

Imagine a universe whose only sentient inhabitants are self-sufficient, divine beings who devote their lives to various solitary intellectual activities, as they judge best, and want nothing from one another. Having no need or desire whose satisfaction requires interaction with other beings, they act in total isolation from one another. They are never tempted to act contrary to their better judgment and they have no frivolous desires. They also have no reactive attitudes: indignation, gratitude, and the like. Nor do they have any concept of such attitudes. Finally, they know nothing of morality and moral reasons for action. Might they sometimes act freely?

One of these beings, Zed, has devoted the past year to working out the details of a variety of possible geometries. Today he is thinking about what task to turn his considerable intellectual powers to next. The candidates he is considering include modal logic, probability theory, and decision theory. In the end, after much rational thought, Zed decides to tackle decision theory next. Might he have made that decision freely? I do not see why not—provided that free decisions are possible. And, given that, through no fault of his own, Zed has no grasp of morality and moral reasons, it is false that he is morally responsible for his decision.

The idea that morality does not apply to Zed's actions is not essential to this story. If someone claims that morality does apply to his conduct because it would be morally impermissible for him to fritter away his time on trivial pursuits, even though he has no idea this is so, I do not object. The same goes for the claim that morality applies to his conduct because some of what he does is morally permissible: that is, I do not object to this claim either. Zed's having, through no fault of his own, no conception of morality and moral reasons is itself plausibly regarded as sufficient for his lacking moral responsibility for his actions.³

5.3. The Hard Question

My fourth question was this: Can agents who never act freely be morally responsible for some of their actions? My discussion of the preceding three questions was guided by a number of cases. With the fourth question, I take a different tack. It includes what one might think of as philosophical psychoanalysis. But I start with something else—some claims about free will and moral responsibility by a well-known neuroscientist.

Michael Gazzaniga asserts that free will essentially involves a ghostly or nonphysical element and “some secret stuff that is YOU” (2011, p. 108). So it is no surprise that, in his view, “free will is a miscast concept, based on social and psychological beliefs . . . that have not been borne out and/or are at odds with modern scientific knowledge about the nature of our universe” (p. 219). Although Gazzaniga regards free will as an illusion, the main thesis of his book is that we are morally responsible agents and that scientific findings do not undermine the latter claim. In Gazzaniga’s view, free will is magical and moral responsibility is not. “The issue isn’t whether or not we are ‘free.’ The issue is that there is no scientific reason not to hold people accountable and responsible” (p. 106). Of course, someone might say that it is a good idea to hold people responsible even if they are not, but Gazzaniga is not advocating that idea.

Some people agree with Gazzaniga about what “free will” means but part ways with him on moral responsibility. Joshua Greene and Jonathan Cohen argue that “new neuroscience will change the law . . . by transforming people’s intuitions about free will and moral responsibility” (2004, p. 1775). In their opinion, “Most people’s view of the mind is implicitly *dualist* and *libertarian*” (p. 1779); and they contend that neuroscience will undermine this view of the mind, “people’s common sense, libertarian conception of free will and the retributivist thinking that depends on it” (p. 1776).

While I am at it, I reproduce a remarkable passage on free will from neuroscientist P. Read Montague.

Free will is the idea that we make choices and have thoughts independent of anything remotely resembling a physical process. Free will is the close cousin to the idea of the soul—the concept that “you,” your thoughts and feelings, derive from an entity that is separate and distinct from the physical mechanisms that make up your body. From this perspective, your choices are not caused by physical events, but instead emerge wholly formed from somewhere indescribable and outside the purview of physical descriptions. This implies that free will cannot have evolved by natural selection, as that would place it directly in a stream of causally connected events. (2008, p. 584)

Someone who holds Montague’s view of free will may or may not see the view as having important implications for moral responsibility. Such a person may side with Gazzaniga in rejecting free will while holding on to moral responsibility or side with Greene and Cohen in rejecting both.

Gazzaniga asserts that having free will depends on having or being an immaterial soul, whereas being morally responsible for actions one performs does not. That is one way of opening the door to the idea that agents who never act freely may be morally responsible for some of their actions. But it is not the only way. For example, someone may claim that the combination of being or having a soul, having agent-causal powers (see chapter 1, section 1.1), and being an indeterministic decision-maker is required for free will but that only two of these things are required for moral responsibility. Someone else who agrees with this claim about free will may claim that only one element of this combination is required for moral responsibility. And another person may contend that although all of these things are required for free will, none are required for moral responsibility. Theoretical options abound.

Peter van Inwagen has said that if science were to “present us with compelling reasons for believing in determinism,” we should “become compatibilists” about free will (1983, p. 223). Consider the following three propositions (van Inwagen, p. 219):

1. “We are sometimes morally responsible for the consequences of our acts.”
2. “If 1 is true, then we have free will.”
3. “Our having free will entails indeterminism.”

Although van Inwagen believes that all three propositions are true, he also believes the following: if either 1 or 2 is false, then 1 is true and 2 is false; and if either 2 or 3 is false, then 2 is true and 3 is false (p. 220). To make things simpler, one can say that regarding these three propositions, van Inwagen is least confident in the truth of 3 and most confident in the truth of 1. That is why he would give up 3 rather than either 1 or 2 if it were shown that determinism is true. (Van Inwagen reports that what I am representing as differences in degrees of confidence are explained by his assessments of the relative strengths of his arguments for these propositions [p. 223].)⁴

Imagine a philosopher, Phyllis, who, like van Inwagen, assents to propositions 1, 2, and 3 but is more confident of 3 than 2. If she were to become convinced that determinism is true, she might give up 2 and hold on to 1 and 3. Phyllis might contend that although moral responsibility is compatible with determinism, free will is not.

It is sometimes said that John Fischer holds the view just mentioned—that his semicompatibilism is compatibilism about moral responsibility and incompatibilism about free will. This misrepresents Fischer’s position. He describes his semicompatibilism as the view that “moral responsibility is compatible

with causal determinism, even if causal determinism is incompatible with freedom to do otherwise" (1994, p. 180). Fischer's semicompatibilism also is a view about free action. He asserts that "guidance control is the freedom-relevant condition necessary and sufficient for moral responsibility" (p. 168), and he reports that his "account of guidance control (and moral responsibility) . . . yields 'semicompatibilism'" (p. 180). Thus, I take semicompatibilism to encompass the thesis that free action is compatible with determinism (as traditional compatibilists assert), even if "determinism is incompatible with freedom to do otherwise" (which traditional compatibilists deny). Someone who believes that free will is the ability to act freely or that having free will does not depend on having the freedom to do otherwise may contend that free will is compatible with determinism even if "determinism is incompatible with freedom to do otherwise." In any case, Fischer identifies "acting freely" with "exercising guidance control" (2006, p. 21), and, in his view, guidance control can be exercised in deterministic worlds. In the same vein, Fischer and Ravizza write: "Guidance control of an action involves an agent's freely performing an action" (1998, p. 31).⁵

A philosopher who does float the idea that even if moral responsibility is compatible with determinism, free will is not is Ted Warfield (2003, p. 621). He floats it in a discussion of a Frankfurt-style case. A brief description of my favorite Frankfurt-style story helps set the stage for my discussion of the idea. Readers who are interested in more detailed descriptions of the story are referred to Mele and Robb 1998 and 2003.

Frank initiates a process *P* in Bob's brain at *t*₁ with the intention of thereby causing Bob to decide at *t*₂ (five minutes later, say) to steal Ann's car. The process, which is screened off from Bob's consciousness, will culminate in Bob's deciding at *t*₂ to steal Ann's car unless he decides on his own at *t*₂ to steal it or is incapable at *t*₂ of making a decision (because, for example, he is dead by *t*₂). Indeed, *P* makes it impossible for Bob to avoid deciding at *t*₂ to steal Ann's car, given that he is capable of making a decision then. *P* operates in such a way that unless *P* is preempted by some event or process external to it, it will cause Bob to decide at *t*₂ to steal Ann's car. And given that Bob is capable of making a decision at the time, the only thing that can preempt *P* is Bob's deciding on his own at *t*₂ to steal Ann's car. Thus, all worlds with the same laws of nature as *W* and the same past up to *t*₂ are worlds in which Bob decides at *t*₂ to steal Ann's car. As it happens, at *t*₂ Bob decides on his own to steal the car, on the basis of his own indeterministic deliberation about whether to steal it. But if he had not just then decided on his own to steal it, *P* would have issued, at *t*₂, in his deciding to steal it.⁶ Rest assured that *P* in

no way influences the indeterministic decision-making process that actually issues in Bob's decision.

On the supposition that an indeterministic event-causal process, x , issues in Bob's decision, that x does not deterministically cause the decision is clear: even though x actually causes Bob's deciding to steal Ann's car, which decision is made at t_2 , there are possible worlds with the same laws of nature as W and the same past up to t_2 in which Bob is capable at t_2 of making a decision, x is not preempted or disturbed in any way by anything external to it, and, even so, x does not cause Bob's deciding at t_2 to steal Ann's car. (In those worlds, P causes Bob's deciding to do that, which decision is made at t_2). P makes it inevitable that Bob will decide at t_2 to steal the car, if at t_2 he is capable of making a decision at t_2 , but since P does not actually cause Bob's deciding to do that, it does not deterministically cause it.

Build it into the story that Bob is sane and rational, understands why stealing the car would be wrong, and so on. Some readers will find it intuitive that Bob is morally responsible for deciding to steal the car. After all, he decided on his own to steal it. P played no role in producing his decision. Some of the same readers will also find it intuitive that Bob freely decided to steal the car—that he made this decision of his own free will. But what about Phyllis? Might she be willing to claim that Bob is morally responsible for his decision but unwilling to say that he freely made it? Recall that Phyllis's confidence in incompatibilism about free will outstrips her confidence that the existence of morally responsible agents depends on the existence of free will. But that fact itself does not directly yield an answer to my question. We have to look deeper, and asking Phyllis what it is about determinism in virtue of which it is incompatible with free will might be a good idea. I asked. And she replied that a necessary condition for an agent's freely deciding to A is that there be another possible world with the same laws as the world at issue and the same past right up to the moment at which the agent decides to A in which the agent does not decide to A . Phyllis observes that this condition is not satisfied in any deterministic world and that it is not satisfied in my story about Bob either, despite that story's having an indeterministic setting. She adds that she is more confident about the truth of this alleged necessary condition than she is about the truth of the claim that the existence of free will is necessary for the existence of moral responsibility.

Phyllis maintains that an agent's decision to A is directly free only if he had the freedom to do otherwise at the time than decide to A . And she is more confident that this is true than she is that an agent is directly morally responsible for deciding to A only if he had the freedom to do otherwise at the time

than decide to *A*. As I mentioned, Phyllis understands the freedom to do otherwise as essentially indeterministic.

Imagine that you agree with Phyllis that some agents in Frankfurt-style cases are directly morally responsible for deciding to *A* even though they could not have done otherwise than decide to *A*. How might you try to persuade her that these agents freely decide to *A*?

You might try to persuade Phyllis that, necessarily, if an agent is directly morally responsible for *A*-ing, then he freely *A*-s. And how might you try to do that? You might start by trying to produce cases in which an agent is not morally responsible for *A*-ing because he *A*-s unfreely. You invite Phyllis to consider a compulsive handwasher who has just now washed his hands for the fiftieth time today, harming himself and his family in the process. And you assert that this unfortunate person is not morally responsible for this action because he unfreely washes his hands.

Phyllis regards your diagnosis of the person's nonresponsibility as undiscerning. She says that he lacks moral responsibility for the handwashing because the action was *compelled*. And she proceeds to distinguish between compulsion and ordinary deterministic causation. One common compatibilist charge against incompatibilists is that they confuse deterministic causation with compulsion.⁷ Stock examples of compulsion are a hypothetical kleptomaniac's stealing something because he has an uncontrollable urge to steal it and a hypothetical addict's using a drug because of an irresistible desire to use it (which is not to say that ordinary kleptomaniacs' desires to steal or ordinary addicts' desires for drugs are often—or ever—actually irresistible). Compatibilists typically hold that actions such as these in deterministic universes are importantly different from, say, an ordinary person's buying a pack of gum or taking an aspirin in a deterministic universe. And Phyllis agrees. In her view, determinism precludes the freedom to do otherwise that is required for directly free action, but it is compatible with direct moral responsibility for actions. Compulsion, she says, is incompatible both with the freedom to do otherwise and with direct moral responsibility.

Assume that Phyllis will not budge on her incompatibilism about free will. You might try to persuade her that moral responsibility also is incompatible with determinism by arguing for the following thesis: Necessarily, if a being performs at least one action for which he is morally responsible, he performs at least one action freely. I leave producing such an argument as an exercise for the reader. Another tack is to argue directly for incompatibilism about moral responsibility. There is, of course, a well-known argument for this thesis, Peter van Inwagen's direct argument (1983, pp. 183–88). That argument has been

in circulation for over thirty years and has attracted a lot of attention. Some philosophers are persuaded by it and others are unpersuaded. Phyllis may fall into the latter group. She may be unpersuaded by the direct argument, even if she finds the consequence argument persuasive.

Is Phyllis just being stubborn? Consider the following claim: (*BI*) Incompatibilism about free will is true if and only if incompatibilism about moral responsibility is true. Is *BI* obviously true? Seemingly not. I leave constructing a knock-down argument for it as an exercise for the reader.

5.4. More on the Hard Question

In the preceding section, I explored the idea that reflection on a Frankfurt-style case might help lead a philosopher to the combination of an incompatibilist view of free action and a compatibilist view of moral responsibility. A philosopher who takes this stand is in a position to hold that there are possible worlds in which some agents are morally responsible for some of their actions even though no agent ever acts freely. My fourth question, again, was this: Can agents who never act freely be morally responsible for some of their actions? Such a philosopher answers *yes*.

In the present section, I explore a route to a *yes* answer that goes through a Frankfurt-style case but embraces incompatibilism both about free action and about moral responsibility. Phyllis took a compatibilist line on moral responsibility in response to an argument from compulsion. Her brother, Phil, who is an incompatibilist both about moral responsibility and about free will, takes a different line. He holds that being compelled to decide to *A* is importantly different from the *inevitability* of an agent's deciding to *A* in an indeterministic Frankfurt-style case of the sort sketched in section 5.3. In his view, as in Phyllis's, that inevitability is incompatible with Bob's freely deciding to *A*. But, while holding on to his incompatibilism about moral responsibility, he maintains that the inevitability at issue is compatible with Bob's being directly morally responsible for deciding to *A*.

Phil takes a *leeway* incompatibilist position on free action and a *source* incompatibilist position on moral responsibility.⁸ Anyone who holds that an agent's decision to *A* in a world *W* is directly free only if there is another world with the same past right up to the time of decision and the same laws of nature in which he does not decide to *A* at that time is, by definition, a leeway incompatibilist about directly free decision. Phil views free action as requiring the freedom to do otherwise, and he interprets that freedom in the incompatibilist way just mentioned. But he is persuaded by a Frankfurt-style case that an

agent can be directly morally responsible for deciding to *A* even though he does not satisfy his leeway incompatibilist condition for free action. (Source incompatibilists about free action are persuaded by Frankfurt-style cases to take a parallel view about directly free decisions: what matters is how the decision was produced, not leeway.)

How do we get from here to the position that agents who *never* act freely can be morally responsible for some of their actions? With the assistance of a “global” Frankfurt story, a story in which, *whenever* an agent acts, he acts “on his own” but could not have done otherwise than perform the action he performed, owing to the presence of a process like *P* in the Mele-Robb story.⁹ Such a story obviously involves an amazing string of coincidences, but that feature does not render the story impossible. People who see Bob as morally responsible for deciding to steal Ann’s car in the Mele-Robb story may see him as morally responsible for various actions he performs in a globalized version of the story. But any leeway incompatibilists about directly free actions among them will see him as never acting freely.

Return to the one-shot Mele-Robb story. Bob decides on his own to steal the car, and he is sane, rational, and reasons-responsive. He views himself as a moral agent and understands why stealing the car would be wrong. On the basis of his own indeterministic deliberation, he decides to steal the car. The following two questions (among others, of course) arise: Was Bob directly morally responsible for deciding to steal the car? And did he directly freely decide to steal the car? Both questions are about the same decision, a decision that came about in the way described and in the circumstances described. If the answer to the first question is yes, then Bob satisfies any “control” requirements for being directly morally responsible for a decision. Some people will claim that satisfaction of these conditions is sufficient for Bob’s decision’s being directly free. A source incompatibilist both about moral responsibility and about free action can claim that in satisfying the control requirements for direct moral responsibility here, Bob satisfies conditions sufficient for his decision’s being directly free. A source incompatibilist about moral responsibility who is a leeway incompatibilist about directly free action may say that the control conditions Bob satisfies are close enough to the control conditions for directly free decisions for Bob to be counted as directly morally responsible for his decision, but fall short of what is needed for directly free action.

Once we get to this point, someone will surely ask why we should believe that the control conditions for directly free action are more demanding than the control conditions for directly morally responsible action. This takes us back to familiar turf—arguments for leeway incompatibilism about directly

free action and for leeway incompatibilism about direct moral responsibility. Someone who is persuaded by some argument for the former but unpersuaded by any argument for the latter may indeed claim that stronger control conditions are required for directly free action than for being directly morally responsible for an action and then cite his or her assessments of these arguments as partial support for the claim.

A parallel point applies to the combination of compatibilism about morally responsible action and incompatibilism about free action. Someone who is persuaded by some argument for the latter (say, some version of the consequence argument) but unpersuaded by any argument for rejecting the former may also claim that more demanding control conditions are required for directly free action than for being directly morally responsible for an action. This also takes us back to familiar turf.

I will not assess and compare the merits of the leading arguments for the following theses: incompatibilism about free will, incompatibilism about moral responsibility, leeway incompatibilism about free will, leeway incompatibilism about moral responsibility. Fortunately, these arguments are very familiar, and perhaps many readers already have definite opinions about their merits. Rather than move onto this familiar turf, I briefly take up a related issue.

5.5. Control and the Importance of Free Will

Suppose that agents who never act freely can nevertheless be morally responsible for some of their actions. That is, suppose, to use some shorthand, that free will is not required for moral responsibility. How important is free will, in that case, and where does its importance lie? When things get complicated in classroom discussions of free will, a bold student will occasionally ask why free will matters anyway. What is so important about it that the professor is willing to examine argument after argument and thought experiment after thought experiment? The reply that free will is required for moral responsibility is a handy motivational tool at that point. The reactive attitudes are pervasive parts of life, and one can easily bring them into contact with moral responsibility. But suppose this tool actually misleads students because free will is not required for moral responsibility. What then?

One can say that free will enters into our self-image in a deep way. One may claim that we see ourselves as having significant control over much of our conduct, including many of our decisions, and that the control we regard ourselves as having is something we indeed have only if we have free will. If this

is right, one may add, free will is important at least partly because our lacking it would involve our being seriously mistaken about what kind of being we are. The discovery that free will is an illusion, one may claim, would be an even greater assault on our self-conception than the discovery that love is an illusion.

Theorists who take this route should look carefully into what sort of control most people see themselves as having over much of their conduct. After all, these theorists are making a claim about control that people regard themselves as having. Does the control most people attribute to themselves require their having some of the things mentioned by the scientists I quoted earlier—“some secret stuff” that is them, souls, or supernatural powers, for example? If not, which of the leading theories of free will best captures the self-attributed control?

Imagine that some compatibilist theory of free will does a good job of capturing the self-attributed control at issue. In that case, the combination of a proof that compatibilism about free will is false and a proof that the self-attributed control is all the control that is needed for moral responsibility would leave lay folk pretty much where they were. The proof of incompatibilism about free will would show that the control that they attribute to themselves and is woven into their self-image falls short of what is needed for free will. But this does not threaten the aspect of their self-image at issue. That is, the imagined proofs do not challenge the proposition that they have the control they attribute to themselves. If, at *this* point, students were to ask why free will matters, the self-image card that I described would be powerless. And some professors might find themselves saying that free will matters in roughly the way that some other things matter that attract little interest outside philosophy—universals and tropes, for example.

Wind the line of reasoning that I have just run through back to the point at which I imagined that a compatibilist theory of free will does a good job of capturing the control featured in our self-image. This time, imagine that some event-causal libertarian view captures the control at issue. Now suppose that we have a proof that although the incompatibilist aspect of this libertarian view is correct, free will requires agent causation. And suppose that we also have a proof that the self-attributed control at issue is all the control that is needed for moral responsibility; agent-causal powers are not needed for this. Here again, lay folk would be left pretty much where they were. The revelation that they lack free will would not have much of an impact. After all, the imagined proofs do not challenge the proposition that they have the control they attribute to themselves. Similar results can be achieved from other

starting points. For example, substitute an agent-causal libertarian view for an event-causal libertarian view in the scenario just sketched and imagine that we have a proof that free will requires something more—namely, that immaterial souls are at work in decision-making.

In Mele 2007a, I voiced my suspicion that a philosopher who holds that agents who never act freely can nevertheless be morally responsible for some of their actions is “attracted to a combination of conceptions [that is, of free will (or free action) on the one hand and moral responsibility on the other] that makes the gap between free agency and morally responsible agency unacceptably large—so large that adopting these conceptions would prove to be unfruitful” (p. 207). One thing I was worried about then, as I recall, is that one might end up with a conception of free will that is of little interest outside philosophy and that free will would go the way of tropes and universals. But we will not be in a position to know whether this worry is warranted until we know more about—among other things—the control that lay folk attribute to themselves and others. This, I believe, would be an excellent topic of investigation for experimental philosophers.

5.6. Conclusion

I repeat the four questions I said I would explore in this chapter:

1. Can agents be morally responsible for *A*-ing even though they did not freely *A*?
2. Can agents freely *A* without being morally responsible for *A*-ing?
3. Can agents who are never morally responsible for anything sometimes act freely?
4. Can agents who never act freely be morally responsible for some of their actions?

In section 5.2, I defended *yes* answers to the first three questions. In sections 5.3 and 5.4, I explored some ways in which a philosopher may come around to the view that question 4 should be answered *yes*. I am open to this view. That is, I believe that serious arguments for it should be taken seriously. I have not offered an argument for it myself. Nor have I offered an argument against it.

In arriving at answers to my first three questions, attention to cases seemed to do the work that needed to be done. Rather than answer question 4, I explored some ways in which one might arrive at an affirmative answer. To be justifiably confident in one's answer one way or the other, one may need to

defend detailed accounts of free will and moral responsibility and pay very close attention to roles played by control in those accounts. That is a project for a book in its own right, and this is not that book. For reasons of the sort touched on in section 5.5, the author of this possible book may also wish to look into how lay folk conceive of the control they have over their conduct.

What have I accomplished in this chapter beyond defending affirmative answers to my first three questions? One thing I hope I have done is to motivate readers who assume that any possible world without free actions is a world in which no agents are morally responsible for any of their actions to wonder how the assumed proposition can be successfully supported. In the past, I myself have made this assumption; but I have no powerful argument for it. Perhaps some readers believe they do have such an argument. If so, I am eager to see it. Another thing I hoped to do was to highlight a potential cost of rejecting the assumed proposition and to suggest a partial approach to exploring that cost. If free will really should be of interest primarily to philosophers, so be it; but I hope we will learn that it should be of much broader interest.

I had other aims too. While pursuing my guiding questions, I managed to introduce a variety of positions on free will, including semicompatibilism, leeway incompatibilism, and source incompatibilism. This helps set the stage for subsequent chapters. I also illustrated a point that readers should keep in mind: Different incompatibilists about free will have different reasons for endorsing incompatibilism. Some regard determinism as precluding the existence of magical powers that they view as necessary for free will. Such incompatibilists see little value in a naturalistic event-causal libertarianism (a position discussed in detail in subsequent chapters); but if they view agent causation as magical, some of them may find it attractive. Others shun magical requirements for free will and take determinism to preclude free will by precluding leeway. And yet others, moved by Frankfurt-style cases, contend that the real problem with determinism is that it does not leave room for the existence of agents with the power to be indeterministic initiators of some of their decisions. Incompatibilists about free will are a diverse group.

In chapter 1, I reported that my interest in free action is in what I call *moral-responsibility-level free action*—“roughly, free action of such a kind that if all the freedom-independent conditions for moral responsibility for a particular action were satisfied without that sufficing for the agent’s being morally responsible for it, the addition of the action’s being free to this set of conditions would entail that he is morally responsible for it” (Mele 2006, p. 17). I stand by that report, but without insisting that every possible world with morally responsible agents is a world with free agents. Even though I do

not insist on this, I treat the proposition that only agents who sometimes act freely are morally responsible for some of their actions as a working assumption in this book.¹⁰

Notes

1. I owe this way of putting things to Josh Gert.
2. Kane's view has changed. See, for example, Kane 2008.
3. For encouraging reactions to Zed's story, I am grateful to Josh Gert, Ish Haji, David McNaughton, and Piers Rawling. For a related story, see Mele 1995a, pp. 3–4.
4. For a critique of van Inwagen's argument for 1, see Mele 1995a, pp. 243–46.
5. Of course, it is open to someone to claim that free action does not depend on free will and that free will requires the freedom to do otherwise.
6. In Mele 2003a (chap. 2), I argued that actions (including decisions) are, essentially, events with a causal history of a certain kind. One might worry that no event that *P* can produce at *t*₂ can be produced by *P* in a way consistent with the event's being a decision. The worry might, for example, derive from the thought that all decisions have beliefs and desires of the agent among their causes and that this would not be true of an alleged decision in which *P* issues. However, a process like *P* may be designed to produce relevant beliefs and desires for use in producing a decision unless it detects that such beliefs and desires are already present; and if suitable beliefs and desires are already present (as they are in Bob's case), *P* can use them in producing the decision.
7. See Audi 1993, chaps. 7 and 10; Ayer 1954; Grünbaum 1971; Mill 1979, chap. 26, esp. pp. 464–67; and Schlick 1962, chap. 7. Also see Hume's remarks on the liberty of spontaneity versus the liberty of indifference (1739, bk. II, pt. III, sec. 2).
8. Neal Tognazzini reports that “the term ‘source incompatibilism’ can be traced back to [McKenna 2001], though the idea had been around for much longer” (2011, p. 75 n. 3). I float a source incompatibilist view in response to Frankfurt-style cases in Mele 1996. (I was then—and still am—officially agnostic about compatibilism, both regarding free will and regarding moral responsibility.)
9. On global Frankfurt-style cases, see Fischer 1994, p. 214; Mele 1995a, p. 141, 1996, pp. 129–39, and 2006, pp. 94–95; and Mele and Robb 1998, pp. 109–10.
10. This chapter is based on Mele 2015a, which is in turn based on a talk I gave in March, 2014 at Queen's College, Oxford University. I am grateful to the audience for discussion.

Is What You Decide Ever up to You?

SOME PHILOSOPHERS HAVE worried about whether libertarians have the resources to explain how agents can have enough control over what they do to perform directly free actions. In Mele 2006, I develop a libertarian response to this worry that is focused on deciding. One plank in the response is the thesis that even if the difference between what an agent does at t in one possible world and what he does at t in another possible world with the same past up to t and the same laws of nature is just a matter of luck, the agent may perform a directly free action at t in both worlds (Mele 2006, chap. 5). I dub this thesis *LDF*.

One may contend, on the following grounds, that *LDF* is false: (1) if a difference of the kind specified in *LDF* is just a matter of luck, then it is not *up to* the agent what he does at t in either world, and (2) if it is not up to one what one does at t , one does not perform a directly free action at t . In section 6.4, I examine an interesting argument against *LDF* along these lines, and I argue that it is unconvincing. Sections 6.1 through 6.3 are written with the dual aim of providing a context for the interesting argument and providing background on what it might be for something to be up to an agent, a topic on which I hope to shed some light. Section 6.5 offers a positive suggestion about its having been up to an agent whether he *A*-ed or *B*-ed, and section 6.6 wraps things up.

I am tempted to say that whether you read on or not is up to you. But I should be more cautious. On some readings of “up to you,” this is true only if you have agent-causal powers (see chapter 1, section 1.1); and there are significant grounds for skepticism about such powers.¹

6.1. Some Background

Brief attention to a theme in a pair of articles by Peter van Inwagen will prove useful as background. A thesis of van Inwagen 2000 is that free will is a mystery even if “agent causation . . . exists” (p. 1). In chapter 4, I discussed the argument he presented there for the thesis that, in a certain indeterministic setting, it is false that he is “able to keep silent” (p. 18) and therefore false that he can freely keep silent. One intriguing feature of van Inwagen’s argument is the link he proposes between being able to *A* and being in a position to make a sincere promise to *A*. I suggested that van Inwagen has a special kind of ability in mind, what I called *promise-level ability*—or *P-ability*, for short. A simplified version of the necessary condition I offered for *P-ability* will do for present purposes:

X is a *P-ability* to *A* only if *X* is a sufficiently reliable ability to ground, in an agent who understands what promising is and knows his own abilities, complete confidence that, barring unexpected excusing factors (including, prominently, unexpected substantial obstacles and unexpected future beliefs that he was tricked into making his promise or mistakenly made it), if he sincerely promises to *A*, he will *A*.

Can an agent who believes that it is now undetermined whether he will *A* or *B* satisfy this condition both with respect to *A*-ing and with respect to *B*-ing? Attention to a story van Inwagen spins in a subsequent article will prove useful in answering this question. How, he asks, will he respond to a certain request to lie to the press if he believes the following:

[*BEL*]. At the moment the question is asked, it will be undetermined whether I shall respond with a lie or with the truth (and therefore it is *now* undetermined which I shall do): if there were a large number of perfect duplicates of me (in identical environments) at the moment the reporter asked her question, some (to a near certainty) would lie and some would tell the truth. And, moreover, it would at the moment the question was asked be undetermined how I should respond to it even if, at some moment between the present moment and that moment, I promised to lie. (2011, p. 478)

He adds: “Perhaps I believe this because I believe that whichever decision I make—to lie or to tell the truth—will be a free decision and believe on philosophical grounds that free decisions must be undetermined events.”

An undetermined decision to *A* might be part of something that “determines” an *A*-ing. If as van Inwagen grants, it might be undetermined whether one will *A* before one promises to *A* and determined that one will *A* after one promises to *A* (2011, pp. 477, 479), it might be undetermined whether one will *A* before one decides to *A* and determined that one will *A* after one decides to *A*; and one’s deciding to *A* may be an important part of what determines that one will *A*. Possibly, in the last sentence quoted above, van Inwagen confuses deciding to do something with doing it. (They are different, of course: for example, one can decide to lie and then fail to lie.) But it is more likely that he is assuming that, in his story, he believes that he will not decide what to do before the question is asked. Notice that, other things being equal, an agent who believes this about himself—while he believes this—is in no position to make at that time a sincere promise to lie. After all, he is unsettled about what to do, he knows that, and he expects his unsettledness to persist until the question is asked. Matters are very different with an agent who has decided to lie and who is fully convinced that, because he has so decided, he will lie. Such an agent would seem to be in an excellent position to make a sincere promise to lie.

A libertarian may claim that an act of lying is not directly free if it is determined by anything at all—even by something that includes an undetermined decision to lie. Suppose this claim is true. An observation about passage *BEL* is in order in this connection. An agent who has the belief described there may also believe—without contradiction, justifiably, and truly—that if he were to decide to lie, it would be undetermined but extremely likely that he would lie. (And, at the same time, he may believe—justifiably and truly—that if he were to decide to tell the truth, it would be undetermined and just as likely that he would tell the truth.) The agent may justifiably believe as well that deciding to lie is open to him, given the past and the laws of nature (and, at the same time, he may justifiably believe the same about deciding to tell the truth). If such an agent decides to lie, he can justifiably take himself to be in a position to make a sincere promise to lie. Suppose he does decide to lie. Then he is entitled to count on himself to lie despite believing, after so deciding, that it is undetermined that he will lie. After all, at that point, he justifiably takes the chance that he will not lie to be minuscule.

An indeterministic agent who is able to raise the probability that he will lie at *t* to nearly 1 just by deciding to lie and able to do the same for the probability that he will instead tell the truth at *t* just by deciding to tell the truth seems pretty impressive—especially if he can do this sort of thing for a considerable range of alternative overt actions and at many times. Setting aside purely stipulative senses of “able,” such an agent, as *t* approaches, is able to lie at

t and able to tell the truth at t , provided that he is able to make each decision. (I am counting on readers not to confuse “each” with “both.”) Furthermore, such an agent may seem to have a lot of control over some of his overt actions. But, one may wonder, how much control does he have over what *decisions* he makes? Van Inwagen’s real worry seems to be about undetermined *decisions*, not about such things as undetermined lying and truth-telling. A worry about such decisions is the topic of section 6.2.

Van Inwagen writes:

[*P*]. If one is, at a certain moment, faced with a choice between doing *A* and doing *B*, it is then up to one whether one will do *A* or *B* *only if* it is then undetermined whether one will do *A* or do *B*—and *necessarily* so. (2011, p. 475)

An implication of *P* is that when agents are faced with choices, it is up to them what they will do only if determinism is false. In what follows, I assume for the sake of argument that this is true.

6.2. A Worry about Luck

The preceding section helps set the stage for one way of presenting a worry about luck. A bit of technical terminology will also help. I reserve the label “basically free action” for any free actions—free *A*-ings—that occur at times at which the past (up to those times) and the laws of nature are consistent with the agent’s not *A*-ing then (see Mele 2006, p. 6). If at t someone performs a basically free action of deciding to *A*, then at t in another possible world with the same past up to t and the same laws of nature he does not decide to *A*. Much of the discussion of free action in this chapter is about basically free action. I should add that I use “deciding [to *A*]” and “choosing [to *A*]” interchangeably. My choice of one or the other term in this chapter is often guided by quotations I am discussing.

What I have called “the problem of present luck” (Mele 2006, p. 66) or, more fully, “the problem of present indeterministic luck” (p. 201) does not immediately leap out at everyone, and I have devoted a lot of ink to making the problem salient (Mele 2005; 2006, pp. 5–9, chap. 3, chap. 5). The next several paragraphs are a relatively brief statement of the problem.²

Consider the following story (from Mele 2006, pp. 73–74). Bob lives in a town in which people make many strange bets, including bets on whether the opening coin toss for football games will occur on time. After Bob agreed to

toss a coin at noon to start a high school football game, Carl, a notorious gambler, offered him \$50 to wait until 12:02 to toss it. Bob was uncertain about what to do, and he was still struggling with his dilemma as noon approached. Although he was tempted by the \$50, he also had moral qualms about helping Carl cheat people out of their money. He judged it best on the whole to do what he agreed to do. Even so, at noon, he decided to toss the coin at 12:02 and to pretend to be searching for it in his pockets in the meantime (decided to *C*, for short).

Bob's decision is basically free and he is basically morally responsible for making it only if there is another possible world with the same past up to noon and the same laws of nature in which, at noon, Bob does not decide to *C*. In some such worlds, Bob decides at noon to toss the coin straightaway. In others, he is still thinking at noon about what to do. There are lots of other candidates for apparent alternative possibilities: at noon, Bob decides to hold on to the coin and to begin singing "Stone Free" straightaway; at noon, Bob decides to start dancing straightaway while holding on to the coin; and so on. In the present theoretical context, candidates for apparent alternative possibilities are genuine possibilities if and only if Bob's doing these things at noon is compatible with the actual world's past up to noon and its laws of nature. The genuine possibilities are, as I put it in a recent article (where I avoid putting things in terms of luck), different possible continuations of a (normally very long) world segment (2013d).

Someone may assert that the relevant worlds diverge as they do at noon because, in these worlds, it is up to Bob what he does at noon and he acts differently at noon in these worlds. But, one may ask, is it any more up to Bob at noon whether, right then, he decides to cheat or instead, for example, decides to flip the coin than it is up to a genuinely random number generator whether the number it outputs at noon is 7, 11, or 13 in a scenario in which it has only these three possible outputs at the time (see Mele 2013d, p. 244)? This is among the questions raised by the problem of present luck. One who poses the problem may hope for a persuasive defense of a plausible answer.

Typical libertarians contend that Bob's being directly morally responsible for deciding to *C* and his directly freely deciding to *C* require that at least one other continuation was possible at noon, a continuation in which Bob does something else at noon. Suppose that another possible continuation was Bob's deciding at noon to toss the coin straightaway; in another possible world with the same past as the actual world up to noon and the same laws of nature, that is what happens. This supposition will be viewed as a double-edged sword by some. A philosopher may believe that having control over whether one *A*-s

or does something else instead is required for directly freely *A*-ing and for being directly morally responsible for *A*-ing and believe that having such control requires that *A*-ing at *t* and doing something else instead at *t* are possible continuations of the past up to *t* for the agent. But the same philosopher may worry that these possible continuations are similar enough to possible continuations for the indeterministic number generator that whatever control the agent may have over whether he *A*-s or does something else instead falls short of what is required for directly free *A*-ing and for direct moral responsibility for *A*-ing.

Consider a fuller version of Bob's story in which although—right up to noon—Bob does his very best to talk himself into doing the right thing and to bring it about that he does not succumb to temptation, he decides at noon to *C*. In another possible world with the same past up to noon and the same laws of nature, Bob's best was good enough: he decides at noon to toss the coin straightaway. That things can turn out so differently at noon (morally or evaluatively speaking) despite the fact that the worlds share the same past up to noon and the same laws of nature will suggest to some readers that Bob lacks sufficient control over whether he makes the bad decision or does something else instead to make that decision freely and to be morally responsible for the decision he actually makes (again, it is the direct versions of free action and moral responsibility that are at issue). After all, in doing his best, Bob did the best he could do to maximize the probability that he would decide to do the right thing, and, even so, he decided to cheat. One may worry that what Bob decides is not sufficiently *up to him* for Bob to be directly morally responsible for making the decision he makes and for it to be a directly free decision.

Given the details of Bob's story, how can Bob have enough control over whether he decides to *C* or does something else instead at noon for his decision to be directly free and for him to be directly morally responsible for it? This is an instance of the central question posed by what I called "the problem of present luck" (2005, p. 411; 2006, p. 66) and what I more recently called "the continuation problem" (2013d).

I am not alone, of course, in seeing present luck as a problem to be dealt with. Timothy O'Connor refers to "a chancy element to choice that cannot be attributed to the person" in a representative event-causal libertarian view, and he deems "the kind of control that is exercised . . . too weak to ground [the agent's] responsibility for which of the causal possibilities is realized" (2000, p. 40). O'Connor contends that typical event-causal libertarian views have the following upshot: "There are objective probabilities corresponding to each of the [possible choices], but within those fixed parameters, which

choice occurs on a given occasion seems, as far as the agent's direct control goes, a matter of chance" (p. xiii; see p. 29). He looks to agent causation for a solution to the problem.³

Suppose that if the pertinent difference at t between a world in which an agent decides at t to A and a world with the same past up to t and the same laws of nature in which he decides at t to B is just a matter of luck, then he does not exercise what might be termed "complete control" over whether he decides at t to A or instead decides at t to B . Even so, if *LDF* is true, these decisions may be directly free. I emphasize that by "complete control" I do not mean "as much control as metaphysically possible." For example, I leave it open that the following conjunction is true: exercising the power of agent causation is required for exercising complete control over whether one decides at t to A or instead decides at t to B , and agent causation is metaphysically impossible.⁴

A novice may claim that because the problem of present luck is generated by a typical libertarian requirement for directly free actions it cannot be a problem for libertarianism. An obvious problem with this claim is that something that someone asserts to be a necessary condition for X can be incompatible with X . Consider, for example, the idea that free will requires determinism, which has had some advocates. If incompatibilists are right, that alleged necessary condition for free will is incompatible with free will. Or consider the claim that possessing the power of agent causation is required for having free will. If agent causation is impossible, and the alleged necessary condition is true, then free will is impossible.⁵

Libertarians have other options, of course, for resisting the claim that choices that are indeterministically caused by their proximal causes—or by their agents—are partly a matter of luck or chance in a way that renders the choices unfree. For example, they can make a case for the view that even if these choices are partly a matter of luck or chance, that is compatible with their being directly free choices (see Mele 2006, chap. 5; O'Connor 2011, p. 325; Steward 2012). And they can say the same about its being up to the agent what he will choose. Another option is to contend that at least some of the choices at issue are not even partly a matter of luck or chance. I return to these options in section 6.3.

6.3. Abilities and "Up To"

The first sentence of van Inwagen 2011 reads as follows: "Let us say that it is at a certain moment *up to* one *whether* one will do A or do B if one is then faced with a choice between doing A and doing B and one is then able to do A and

is then able to do *B*" (p. 475). Of course, to understand the sufficient condition he is proposing here one needs to understand what he means by "able" in this sentence.

Philosophers argue about whether determinism is compatible with an agent's having been able to do otherwise than he did. On one view of the matter, there may be different senses of "able," including a sense in which determinism is compatible with this and another sense in which it is not. On this view, it may be debated whether free action depends on abilities that are precluded by determinism, but the debate is not about the one and only meaning of "able"; it is allowed that "able" may be polysemous.⁶

According to an incompatibilist view of an agent's having been able to do otherwise than he did, an agent who did not do *A* at *t* was able at the pertinent time to *A* at *t* only if in a possible world with the same laws of nature and the same past up to *t*, he *A*-s at *t*.⁷ This view is about what I dub being *O*-able to *A*. On this view, as I pointed out in chapter 4, if agents in deterministic worlds are able to do anything at all, they are able to do only what they actually do. For in any world with the same past and laws as *S*'s deterministic world, *W*, *S* behaves exactly as he does in *W*. In deterministic worlds, agents have no *O*-abilities.

Philosophers who are broad-minded about ability can distinguish *S* and *I* varieties of *O*-ability, to use some shorthand from chapter 4. They can say that an agent in a world *W* who did not do *A* at *t* had, at the relevant time, the *simple O*-ability to *A* at *t* if and only if there is a possible world with the same past and laws as *W* in which he *A*-s at *t*. And they can say that an agent in a world *W* who did not do *A* intentionally at *t* had, at the relevant time, the *O*-ability to *A* intentionally at *t* if and only if there is a possible world with the same past and laws as *W* in which he *A*-s intentionally at *t*.

O-ability to *A* intentionally obviously is stronger than simple *O*-ability, but it is not strong enough for van Inwagen's purposes. Consider a 90% free throw shooter who unintentionally failed to sink his last free throw attempt. The ball hit the front rim, bounced on it a couple of times, and fell off. In another possible world with the same laws and with the same past up to the time at which the ball hits the rim, the ball bounces on the rim a couple of times and into the hoop. Most people would say that, in that world, the player's sinking the free throw is an intentional action. Even so, given his success rate, he would seem not to have been in a position sincerely to promise to sink the shot.

Recall the necessary condition for *promise-level ability* (*P*-ability) stated earlier: *X* is a *P*-ability to *A* only if *X* is a sufficiently reliable ability to ground,

in an agent who understands what promising is and knows his own abilities, complete confidence that, barring unexpected excusing factors (including, prominently, unexpected substantial obstacles and unexpected future beliefs that he was tricked into making his promise or mistakenly made it), if he sincerely promises to *A*, he will *A*. In section 6.1, I explained how an agent who believes that it is now undetermined whether he will *A* or *B* can satisfy this condition for *P*-ability both with respect to *A*-ing and with respect to *B*-ing. An agent with this belief can have promise-level confidence that if he chooses to *A* he will *A* and that if he chooses to *B* he will *B*, and he can be just as confident that he will make a choice (rather than dithering and failing to make one).

The necessary condition just stated for *P*-ability is silent on the question whether an agent's having been able to do otherwise than he did is compatible with determinism. Of course, an incompatibilist necessary condition for *P*-ability can be added. But even an agent's having indeterministic *P*-abilities to keep silent and to spill the beans may be viewed as falling short of what he needs if it is to be *up to him* whether he keeps silent or reveals the damaging fact. It may be claimed that it needs to be up to him what he *chooses*.

Consider a pair of worlds (W_1 and W_2) with the same past up to t and the same laws of nature. In W_1 at t Peter chooses to divulge a damaging fact about a friend; and in W_2 at t he chooses to keep silent about that fact. Now, given that a person's choosing to *A* is itself an intentional action, in W_1 Peter had the *O*-ability to do something intentionally at t that he did not do then—namely, choose to keep silent. And we can say that he had in that world the following pair of abilities: the ability to choose at t to keep silent and the ability to choose at t to spill the beans. (We can say the same about Peter in W_2 .) Each of the abilities at issue is both an *O*-ability and an *I*-ability. Is this enough for it to be up to Peter what he chooses?

Return to the quotation with which I opened this section: "Let us say that it is at a certain moment *up to one whether* one will do *A* or do *B* if one is then faced with a choice between doing *A* and doing *B* and one is then able to do *A* and is then able to do *B*" (van Inwagen, 2011 p. 475). One might try to generate a proposed sufficient condition for its being up to one at a moment whether one will choose to do *A* or choose to do *B* simply by substituting for "do *A*" and "do *B*" in the quotation "choose to do *A*" and "choose to do *B*" and by replacing "doing *A*" and "doing *B*" with "choosing to do *A*" and "choosing to do *B*." But the result would be awkward. (Try it and see. How comfortable are you with the idea of having a choice between choosing to do *A* and

choosing to do *B*?) Using the quotation as a partial model, one might try the following instead:

C. It is at a certain moment up to one whether one will choose to do *A* or choose to do *B* if one is then faced with a choice between doing *A* and doing *B* and one is then able to choose to do *A* and is then able to choose to do *B*.

And one might make it explicit that an agent's having the pair of abilities mentioned in *C* is understood to depend on its being undetermined at the time whether he will choose to do *A* or choose to do *B*.

If "able" in *C* is read as "*I*-able," then Peter in my story satisfies *C*. He is *I*-able to choose to keep silent and *I*-able to choose to divulge the damaging fact, and at no time is it determined which choice he will make. However, just as *I*-ability is not in general sufficient for *P*-ability (or the promise-level ability van Inwagen had in mind, if that differs from *P*-ability), Peter's dual *I*-ability regarding his candidates for choice will not satisfy van Inwagen (see 2000, p. 17). That dual *I*-ability is compatible with the chanciness that worries him.

Some have claimed that agent causation solves the problem about chanciness at the time a choice is made. Rebutting that claim is one of the main purposes of van Inwagen 2000. He contends that "the concept of agent causation is of no use to the philosopher who wants to maintain that free will and indeterminism are compatible (p. 1) and, indeed, "is entirely irrelevant to the problem of free will" (p. 11). Van Inwagen's old promising argument is supposed to help show that this is so (2000, pp. 17–18).

Partly in response to van Inwagen 2000, O'Connor reports that "The agent causationist takes agential control of a freedom-grounding sort as a primitive, both ontologically and conceptually. She then tries to motivate this posit by showing how one might integrate such a primitive feature of control within a wider system of concepts concerning causation, properties, guidance by reasons and so forth" (2011, p. 324). Now, suppose an event-causal libertarian were to say that Peter, in my story, has "agential control of a freedom-grounding sort" over what he chooses in virtue of his being rational, his being well-informed about the pros and cons of his options, his having the dual *I*-abilities I mentioned, and there being no point at which it is determined what he will choose. An agent causationist may say that, even then, it is partly a matter of luck what Peter chooses and therefore is not up to him what he chooses (see O'Connor 2011, p. 324).

Suppose now that Peter has agent-causal powers. In scenarios of the sort at issue, it was not determined that Peter would make the choice he made. In another possible world with the same past up to the moment of decision and the same laws, he makes the opposite choice. So is the difference that I just now mentioned between the two worlds just a matter of chance (or luck)? And is it partly a matter of chance (or luck), therefore, that Peter decides or chooses as he does?

6.4. Do Luck and Free Decisions Mix?

Agent causationists who feel inclined to answer the pair of questions just raised have options. They can answer yes to both, and they can add that the chanciness or luck involved does not preclude the agent's freely making the decision he makes (see O'Connor 2011, p. 325) because agent-causal powers are control powers of a freedom-grounding sort. Or they can answer no to both, and they can add that if the difference at issue were a matter of chance or luck, the agent's decision would not be free. (I am not claiming that these two answers exhaust the possibilities.)

Randolph Clarke has argued that the first option is not a genuine one, at least when it is formulated in terms of luck. His argument is interesting, and attention to it will prove useful. To save ink in the quotation of it that follows, I use some abbreviations: "DT[-s]" replaces "decide[s] to tell the truth" and "DL[-s]" replaces "decide[s] to lie."

Fred freely does something at t [namely, DT] such that, were he to do it, it *would* be the case that at t he DT-s rather than DL-s. Fred is thus able so to act. And Fred is able to do something at t [namely, DL] such that, were he to do it, he would do it freely, and it would then *not* be the case that at t he DT-s rather than DL-s. Then, the fact that at t he DT-s rather than DL-s depends on which of the things Fred is able to do at t he in fact freely does then. It is no stretch of the imagination to suppose that Fred is aware of this dependence. It would seem, then, that it is up to Fred whether, at t , he DT-s rather than DL-s. If that contrastive fact is up to Fred, then it is not just a matter of luck that at t Fred DT-s rather than DL-s. And if this contrastive fact is not just a matter of luck, neither is the difference between Fred's deciding at t to tell the truth and his deciding at t to lie, nor is the difference at t between the actual world, where Fred DT-s, and $\mathcal{W}$, where he DL-s. It thus seems that one cannot consistently accept the luck claim [namely,

that the difference at t between the actual world, where Fred DT-s, and world W , where he DL-s, is just a matter of luck] and hold that Fred's actual decision is free and his alternative decision would have been, too. (Clarke 2004, p. 58)⁸

Is this argument sound? Clarke observes that his concern is with “*directly* free actions” (2004, p. 47)—as he puts it, free actions that do not derive their freedom from “any earlier action” the agent performed (2003, p. 63).⁹ I start with the part of the argument that is supposed to justify the first “up to Fred” claim. If Clarke’s argument is about directly free actions in general, that part of the argument appears to invoke the following principle:

UT. If the fact that at t an agent S A -s rather than B -s depends on which of the things he is able to do at t he in fact (directly) freely does at t and he is aware of this dependence, then it is up to S whether at t he A -s rather than B -s.

This is an interesting principle, and an examination of it will prove instructive even if Clarke did not intend to rely on it. Possibly, he meant his argument to apply only to directly free decisions. I take up that possibility later.

Is *UT* true? Consider the following story.¹⁰ Bart and his neuroscientist friends are playing a betting game. Bets are placed on whether a player will raise his right index finger or his left index finger next. The players’ hands are placed palm down on a table, and the finger raisings are to be performed while their palms remain on the table. The players know that a device is present that will paralyze one finger or the other as soon as a player starts trying to raise either. In fact, each player has agreed to try to raise his right index finger while also trying to raise his left index finger, and it is undetermined which finger the machine will paralyze when the simultaneous dual tryings start.

It is Bart’s turn to raise a finger. He simultaneously tries to raise his right index finger and tries to raise his left index finger. His primary goal is to raise a finger: either will do. He succeeds in raising his right index finger. His left index finger became paralyzed as soon as his attempts began.

Just as Clarke supposes that Fred freely decides to tell the truth, I suppose that Bart freely raises his right index finger. The fact that at t Bart raises his right index finger rather than his left index finger depends on which of the things Bart is able (before paralysis sets in) to do at t that he in fact freely does then. It is no stretch of the imagination to suppose that Bart is aware of this dependence, and I stipulate that he is aware of it. The conjunction of *UT* and

the details of the case entails that it is up to Bart whether, at t , he raises his right index finger rather than his left index finger. (Notice *UT*'s "able to do at t ," which is distinct both from "able throughout t to do at t " and from "able at t to do at t ." See note 14.) But, obviously, this is not up to Bart once he starts trying to raise each finger. (Earlier, it might have been up to him whether to try to raise just one finger.) So either *UT* is false or there is something wrong with my story, which includes the supposition that Bart's raising his right index finger is a free action.

Am I perhaps using "able" in a way that departs from Clarke's usage of the term in his argument? I repeat the first two sentences of the long quotation from Clarke 2004: "Fred freely does something at t [namely, DT] such that, were he to do it, it *would* be the case that at t he DT-s rather than DL-s. Fred is thus able so to act" (p. 58). If the inference here is valid, then so is the following inference: Bart freely does something at t —namely, raise his right index finger—such that, were he to do it, it *would* be the case that at t he raises his right index finger rather than his left; so Bart is able so to act. And Bart is able (up to the moment when paralysis sets in) to do something at t —namely, raise his left index finger—such that, were he to do it, he would (by hypothesis) do it freely, and it would then *not* be the case that at t he raises his right index finger rather than his left index finger. From the perspective of Clarke's inference about what Fred is able to do, if there is a problem with my story, it seems not to be my use of "able."

Might it be that although Clarke is entitled to suppose that Fred freely decides to tell the truth, I am not entitled to suppose that Bart freely raises his right index finger? Again, as Clarke says, it is "*directly* free actions" that are at issue (2004, p. 47), and he characterizes them as free actions that do not derive their freedom from "any earlier action" the agent performed (2003, p. 63). Someone may claim that directly free actions are limited to decisions (or choices), and that is why it is false that Bart directly freely raised his right index finger. Any argument offered for that claim can be assessed. For the record, Clarke himself is opposed to this restrictive idea about directly free actions (2003, pp. 121–26). He asserts that we can acknowledge the importance of decision-making "and still recognize that an action can be directly free even if it is not itself a decision, does not include a decision, and does not result in any direct manner from a decision—indeed, even if the intention-acquisition from which it directly results is not an action at all" (2003, p. 126).¹¹

If there are free actions that do not derive their freedom from "any earlier action" (Clarke 2003, p. 63) and if some overt actions are among them,

I see no good reason to believe that Bart's raising his right index finger cannot also be among them.¹² (I am assuming that nonmoral actions can be free. Readers who reject that assumption should imagine that the game has some minor moral significance.) Nor do I see any good reason to believe that my story about Bart is incoherent. I conclude that *UT* is false and that Clarke's argument does not show that the following is impossible: the difference at *t* between the actual world, where an agent *A*-s at *t*, and world *W*, where he *B*-s at *t*, is just a matter of luck (or chance); and, even so, he directly freely *A*-s at *t* in the actual world and directly freely *B*-s at *t* in *W*.¹³ Let *t* be a stretch of time that begins when Bart simultaneously begins his dual attempts and ends when his right index finger finishes rising. World *W* does not diverge from the actual world until Bart's dual attempts have begun, and Bart has no control over which finger becomes paralyzed. The difference at issue between the actual world and *W* at the time of initial divergence is just a matter of luck or chance.¹⁴

Consider the following restricted replacement for *UT*:

UTD. If the fact that at *t* an agent *S* decides to *A* rather than decides to *B* depends on which of the things he is able to do at *t* he in fact (directly) freely does at *t* and he is aware of this dependence, then it is up to *S* whether at *t* he decides to *A* rather than decides to *B*.

Possibly, Clarke had something like this in mind (and not *UT*). Obviously *UTD* cannot be falsified by my story about Bart. That story is about overt actions—not decisions. Is *UTD* true even though *UT* is false? Are standards for its being up to one whether one decides to do one thing rather than deciding to do another significantly different from standards for its being up to one whether one raises one's right index finger rather than raising one's left? These questions merit some attention.

Clarke writes:

Might one reject the claim that it is up to Fred whether, at *t*, he decides to tell the truth rather than deciding to lie? Whether this contrastive fact holds depends on which of the decisions Fred is able to make he actually makes. We have supposed that whichever of the two decisions Fred makes, he freely makes that decision; and we have supposed that Fred is aware of all this. I do not see how rejecting the claim in question can be consistent with these suppositions, given the dependence. (2004, pp. 58–59)

I have explained how the falsity of a parallel “up to” claim about Bart’s raising his right index finger is compatible with the truth of parallel ability, freedom, dependence, and awareness suppositions about Bart. If it is “up to Fred whether, at t , he [DT-s] rather than [DL-s]” even though it is not up to Bart in my story whether he raises his right or his left index finger, the crucial difference in the two scenarios must lie in one or more differences between deciding and finger raising. And there are some notable differences, including the following two. Finger raisings are overt actions; decidings are not. And decidings (as I conceive of them, at any rate) are momentary actions, whereas finger raisings are not.

Consider the hypothesis that even though it was not up to Bart which of his index fingers he would raise, it was up to Fred right up to t which decision he would make at t . A supporter of this hypothesis may claim that there is a difference between deciding and finger raising that helps to account for its truth—that there is something special about deciding. Might the special thing (if there is one) be that deciding is such that even if the difference at t between a world in which Fred decides at t to tell the truth and a world with the same past up to t and the same laws in which he decides at t to lie is just a matter of luck, it was up to Fred which decision he would make?

If we had an acceptable analysis of “it is up to S whether he decides to x or decides to y ,” we could apply it to the case of Fred with a view to answering the question I just raised. If there is such an analysis, I am unaware of it. In any case, Clarke’s argument does not show that deciding is not special in the way just described. Consequently, even if one will decide freely only if it is up to one what one will decide, Clarke’s argument does not show that a thesis it was designed to falsify is false: namely, the thesis (*LDFd*) that even if the difference between what an agent decides at t in one possible world and what he decides at t in another possible world with the same past up to t and the same laws of nature is just a matter of luck, the agent may make a directly free decision at t in both worlds.

I said that Clarke’s argument does not show that deciding lacks a certain kind of specialness. Why did I say that? Clarke’s argument includes the premise that if “it is up to Fred whether, at t , he [DT-s] rather than [DL-s] . . . then it is not just a matter of luck that at t Fred [DT-s] rather than [DL-s]” (2004, p. 58). But he offers no argument for this premise. Because he does not argue for it and because it is not obviously true when taken at face value, I view Clarke’s assertion as in part a report on how he chooses to use “up to [an agent].” Assertions that are obviously true need no argument, and one may choose to understand the figure of speech at issue in such a way that it

is obvious that the quoted assertion is true. I myself do not regard the assertion as obviously true. (Keep in mind that the premise's consequent is not the assertion that Fred's DT-ing is just a matter of luck and is the assertion that a certain cross-world difference at t is just a matter of luck.) And it merits mention that one who takes one's lead in interpreting "up to Fred" from the grounds Clarke offers in the long passage quoted above for the contention that "It would seem [to be] up to Fred whether, at t , he DT-s rather than DL-s" may reasonably treat it as an open question whether this may be up to Fred even if the difference at t between a world in which he DT-s then and a world with the same past and laws in which he DL-s then is just a matter of luck.

An apparent difference in how we understand the expression "just a matter of luck" may also be relevant. Commenting on the difference between his sinking a basketball shot and his missing it, Clarke writes: "to the extent that I exercise any skill at all, this difference is not just a matter of luck" (2011, pp. 338–39). Readers will recall my basketball example (in section 6.3) in which the featured "sink" and "miss" worlds do not diverge until after the ball leaves the shooter's hands. (They do not diverge until the ball first hits the rim.) In both worlds, the ball bounces on the rim a couple of times. In one it falls into the hoop, and in the other it falls away. The shooter definitely exercises some skill, and I say that the difference at issue is just a matter of luck. Obviously, the difference in outcome is not due in any way to an intrinsic difference in exercises of skills; the worlds do not diverge until after those exercises have ended. It is difficult to see how there can be a substantive disagreement between Clarke and me on this point. I see the apparent disagreement as simply a difference in usage.

Two observations are in order here. First, if the cross-world difference at t in what Fred decides is not "just a matter of luck" in the same sense of the quoted phrase in which the difference in outcome in my basketball example is not "just a matter of luck," then the dispute between Clarke and me about the following thesis is a merely verbal one: Even if the difference at t in what an agent decides in a pair of worlds with the same past up to t and the same laws of nature is just a matter of luck, the agent may freely decide what he decides. For there is an utterly respectable sense of "just a matter of luck" in which the difference in outcome in my basketball example obviously is just a matter of luck. Second, if the operative reading of "just a matter of luck" in Clarke's undefended assertion that if "it is up to Fred whether, at t , he [DT-s] rather than [DL-s] . . . then it is not just a matter of luck that at t Fred [DT-s] rather than [DL-s]" (2004, p. 58) is the same reading at work in the assertion that the difference in outcome in my basketball example is not just a matter

of luck, then I might suppose, for the sake of argument, that the undefended assertion at issue is true. But this supposition leaves it open that on some alternative reading of the expression at issue on which the difference in outcome in my basketball example *is* just a matter of luck, the difference in what Fred decides at t is just a matter of luck even though it is up to Fred at t what he decides then.

Although I know of no acceptable analysis of “it is up to S whether he decides (chooses) to x or decides (chooses) to y ,” some candidates for sufficient conditions have been mentioned here, and so has an alleged necessary condition. Additional attention to them may shed some light on the merits of Clarke’s undefended assertion.

Here are two candidates for sufficient conditions mentioned earlier:

C. It is at a certain moment up to one whether one will choose to do A or choose to do B if one is then faced with a choice between doing A and doing B and one is then able to choose to do A and is then able to choose to do B .

UTD. If the fact that at t an agent S decides to A rather than decides to B depends on which of the things he is able to do at t he in fact (directly) freely does at t and he is aware of this dependence, then it is up to S whether at t he decides to A rather than decides to B .

And here is a candidate for a necessary condition:

UI. Necessarily, if it is up to one right up to t whether one will decide at t to A or decide at t to B , then right up to t it is undetermined whether one will decide at t to A or decide at t to B .¹⁵

Suppose that at time t_0 Fred sets himself the goal of coming to a decision about whether to lie or tell the truth. He then starts rehearsing and weighing pros and cons for each with a view to achieving his goal. His mind is occupied with that task from shortly after t_0 until, a few minutes later, at noon, he decides to tell the truth. In another possible world where everything is the same right up to noon, he decides to lie. Right up to noon, it is undetermined whether Fred will decide to tell the truth (DT) or decide to lie (DL). So Fred’s case does not run afoul of *UI*. And the pertinent instance of *C*’s antecedent is true in Fred’s case when “able” is read as *I*-able. Furthermore, provided that its being undetermined right up to t whether at t Fred will DT or DL is compatible with its being true that the fact that at t he DT-s rather than DL-s depends

on which of the things he is able to do at t he in fact (directly) freely does at t and with his awareness of this dependence, the satisfaction of the pertinent instance of the antecedent of *UTD* can be built into Fred's case. All this leaves it open that the following assertion is true: Even if the difference at t between the actual world in which Fred decides at t to tell the truth and a world with the same past up to t and the same laws in which he decides at t to lie is just a matter of luck, it was up to Fred which decision (or choice) he would make. (As I understand " p leaves it open that q ," it is synonymous with "the falsity of q is not derivable from the truth of p .")

This observation takes us back to van Inwagen's promising arguments (2000, 2011). Consider the following three propositions:

U₁. Necessarily, if it is up to one right up to t whether one will decide at t to A or decide at t to B , then right up to t it is undetermined whether one will decide at t to A or decide at t to B .

U₂. Necessarily, if right up to t it was undetermined whether one would decide at t to A or decide at t to B and one decided one way or the other at t , what one decided was too much a matter of chance or luck for it to have been up to one what one decided.

U₃. Necessarily, one decides directly freely only if it is up to one what one will decide right up to the time at which one decides.

Van Inwagen seems inclined to accept all three propositions. If all three are true, we never make directly free decisions. If we sometimes do make such decisions, at least one of the three propositions at issue is false. A philosopher who believes that we do sometimes decide directly freely and cannot see why any of these three propositions is false may find himself believing, as van Inwagen reports he does (2000), that free will is a mystery.

None of the three propositions at issue is unassailable, of course. Any compatibilist who accepts *U₃* will reject *U₁*. After all, compatibilists reject the idea that directly free decisions depend on the falsity of determinism. A compatibilist proponent of *U₃* will seek to motivate a reading of *U₁*'s antecedent that makes the truth of that antecedent compatible with the truth of determinism.

Some theorists may reject *U₂* on the grounds that some actual or possible beings with the power of agent causation falsify it (O'Connor 2000; Pereboom 2001). Others may argue that what looks to some like chance or luck really is not. And an event-causal libertarian who accepts *U₁* and *U₃* may contend that the chance or luck at issue in *U₂* does not preclude its having been up to one what one decided.¹⁶

What about U_3 ? Any compatibilist who accepts U_1 will reject U_3 . But even libertarians can have doubts about U_3 . Consider a libertarian who, like van Inwagen, conceives of its being up to one what one does in terms of promise-level ability. Such a libertarian may offer an account of promise-level ability to perform overt actions—perhaps one that features objective conditional probabilities of A -ing given that one promises (or decides) to A —and then wonder how promise-level ability to *decide* to A might be understood. Proposition *PD* below should have a familiar ring: it is a candidate for a necessary condition for being a P -ability to A mentioned above, but recast as a necessary condition for being specifically a P -ability to *decide* to A .

PD. X is a P -ability to decide to A only if X is a sufficiently reliable ability to ground, in an agent who understands what promising is and knows his own abilities, complete confidence that, barring unexpected excusing factors, if he sincerely promises to decide to A , he will decide to A .

PD is a strange claim. Have you ever heard anyone promise to *decide* to tell the truth about something—or promise to decide to do anything at all for that matter? I doubt it. And you might be a libertarian who, like van Inwagen, thinks of something's being “up to one” in terms of promise-level ability. If so, your noticing that people never promise to decide to do things (even if they sometimes promise to decide by a certain time *whether* they will or will not do A) may lead you to suspect that talk of its being up to us what we will *decide* is misleading. If it is indeed misleading, then perhaps making directly free decisions does not depend on its being up to us what we will decide—perhaps U_3 is false. (And perhaps not; see below.)

Why don't people ever promise to decide to A ? Perhaps because they realize that sincere promising of this kind would require intending to decide to A . In chapter 2, I discussed the problematic nature of such intending.

Some philosophers will appeal to Frankfurt-style cases as grounds for rejecting U_3 . Return to the story about Bob and the car in chapter 5. Recall that at t_2 Bob decides on his own to steal Ann's car, on the basis of his own indeterministic deliberation about whether to steal it, and if he had not just then decided on his own to steal the car, a certain fail-safe process would have issued, at t_2 , in his deciding to steal it. If the story hits its mark, then even though Bob's world is indeterministic, there is no possible world with the same laws as Bob's world and the same past all the way up to t_2 in which Bob does not decide at t_2 to steal the car. So if its being up to Bob what he will

decide at t_2 depends on there being a possible world with the features just mentioned, it is not up to Bob what he will decide at t_2 . Even so, some philosophers contend, Bob directly freely decides at t_2 to steal Ann's car.

Is there some sense of "up to us" in which it sometimes is up to us what we decide, and might it be up to Fred, in that sense, whether he will DT or DL? I take up this question shortly. First, it should be pointed out that if "up to one" does not have the same meaning in all three of U_1 , U_2 , and U_3 , that trio of propositions does not pose the problem it may seem to pose.

Consider a traditional compatibilist who says he or she endorses U_3 . Such a compatibilist may propose that the following is sufficient for its being up to S what he will decide at t —or, more specifically, for its being up to S whether he will decide at t to A or instead decide at t to B : S is free from compulsion and coercion, is unmanipulated and well informed, has good reasons to A and good reasons to B , is unsettled right up to t about whether to A or B , and, for the duration of his unsettledness about this, is able (on a compatibilist reading of "able," of course) to decide at t to A for reasons that recommend his A -ing and able to decide instead at t to B for reasons that recommend his B -ing. (Such a compatibilist obviously rejects U_1 .) If both compatibilism and U_3 are true, this proposed sufficient condition is attractive.¹⁷

It is open to a libertarian to accept a version of this proposed sufficient condition that differs from it only in that "able" is read as *O*-able.¹⁸ (Recall that *O*-ability is essentially indeterministic.) And accepting it would not force a libertarian to hold that it is up to Bart whether he raises his right or left index finger. Perhaps Bart is free from compulsion and coercion, is unmanipulated and well informed, and has good reasons to raise his right index finger and good reasons to raise his left index finger. But Bart is not unsettled (at the pertinent time) about whether to raise his right or his left index finger; nor, more generally, is he unsettled about what to do at the time. (He is uncertain about which finger he will succeed in raising, but that is another matter.)

Given the proposal at issue, it may be up to Fred whether he will DT or DL even if the pertinent difference at t between a world in which he DT-s at t and a world with the same laws and past in which he DL-s at t is just a matter of chance or luck. In the absence of a convincing argument that this is impossible, the idea merits consideration. Suppose someone were to demonstrate that this idea must be rejected—perhaps because its being up to an agent what he will decide requires something impossible. Would we have to conclude straightaway that no one ever makes directly free decisions? No. One option to explore in light of the imagined newfound understanding of

its being up to an agent what he will decide, is that directly freely deciding to *A* does not require that it was up to one what one would decide. In light of the observation I made in this section about Frankfurt-style cases, this option is considerably less strange than it may *sound*.

Despite what I have said about the assailability of *U*₁, *U*₂, and *U*₃, I believe that they point the way to an interesting question for libertarians. Why aren't cross-world differences of the kind just mentioned at the time of decision incompatible with making decisions that are directly free? Again, if some philosophers are right, a successful answer must feature agent causation (O'Connor 2000). But libertarians who are skeptical about the possibility or existence of agent-causal powers or who do not see how invoking agent causation can solve the apparent problem will hope that these philosophers are wrong (see work cited in note 1). I have offered an alternative solution elsewhere (Mele 2006, chap. 5), and I revisit it in chapter 10. One plank in the solution is an idea—*LDFd*—that the argument by Clarke examined in this section is designed to falsify, and I have shown that his argument is unpersuasive.

6.5. Something Positive

Despite the failure of Clarke's argument to falsify *LDFd*, it usefully brings a pair of questions into close contact. What is it for it to be up to me whether I will decide to *A* or decide to *B*? And what is it to decide freely to *A* (where the freedom at issue is direct)? Both questions require more attention than I have given them here. Although I cannot entirely rectify this shortcoming, an additional suggestion is worth developing briefly.

As I explained in section 6.4, it is not up to Bart whether he will raise his right index finger or instead raise his left index finger. Notice that when Bart freely raises his right index finger, he does not freely omit to raise his left index finger. After all, he is trying to raise his left index finger, and he fails to raise it because the device paralyzes that finger. These points suggest that the following idea is worth considering: (*UTn*) In the case of any uncompelled, uncoerced, unmanipulated, and well informed agent who *A*-ed at *t*, it was up to him whether he *A*-ed then rather than *B*-ed then only if at *t* he freely *A*-ed and freely omitted to *B*.

Consider an agent who was, in van Inwagen's words, "faced with a choice between doing *A* and doing *B*" (2011, p. 475). Might it be that some such agent—an agent who was actively considering pursuing each option—directly freely decided to *A* and, at that time, directly freely omitted to decide to *B*?

And might freely deciding to *A*, in circumstances such as these, be sufficient for freely omitting to decide to *B*? Given a model of free decision-making that features dual attempts to decide—attempting to decide to *A* while also attempting to decide to *B* (see Kane 1999b)—the answer to both questions would seem to be *no*. After all, on this model, the agent who decided to *A* was also simultaneously trying to decide to *B*; and the latter attempt failed. But a dual attempts model of directly free deciding is not the only imaginable model. (For more on this, see chapter 10.)

In section 6.4, I floated the idea that the following is sufficient for its being up to *S* whether he will decide at *t* to *A* or instead decide at *t* to *B*: *S* is free from compulsion and coercion, is unmanipulated and well informed, has good reasons to *A* and good reasons to *B*, is unsettled right up to *t* about whether to *A* or *B*, and, for the duration of his unsettledness about this, is *O*-able to decide at *t* to *A* for reasons that recommend his *A*-ing and *O*-able to decide instead at *t* to *B* for reasons that recommend his *B*-ing. (Again, *O*-ability is essentially indeterministic.) If an agent directly freely decides to *A* while satisfying these conditions, then, in so doing, he directly freely decides *against* doing *B*; and (barring split brains whose separate sides can make conflicting decisions) directly freely deciding at *t* against doing *B* suffices for directly freely omitting at *t* to decide to *B*. So the agent at issue satisfies the condition stated in *UTn*. And in the absence of a convincing argument against *LDFd*, we should take seriously the thought that it was up to this agent which decision he would make, even if the difference at *t* between a world in which he decides at *t* to *A* and a world with the same past up to *t* and the same laws in which he decides at *t* to *B* is just a matter of luck. Some philosophers assume that close attention to something expressed by the idiom “up to the agent” will shed light on deciding freely. Readers will have noticed that the main idea developed in this section about interpreting the idiom is based on attention to deciding freely.¹⁹

6.6. More on *UTn*

A brief commentary on *UTn* is in order before I wrap things up.

UTn. In the case of any uncompelled, uncoerced, unmanipulated, and well informed agent who *A*-ed at *t*, it was up to him whether he *A*-ed then rather than *B*-ed then only if at *t* he freely *A*-ed and freely omitted to *B*.

In the following passage, Carl Ginet seems to *identify* (1) a person's having free will with (2) its being the case that some of his decisions are up to him at the time he makes them: "Kane and I certainly agree that . . . the question whether we have free will—whether any of our decisions are up to us at the time we make them—is . . . an as yet unsettled empirical question about a contingent matter of fact" (2014, p. 25). An alternative view is that although 2 is necessary for 1, 2 is not sufficient for 1. Both views have more specific counterparts. One can identify an agent's freely deciding to *A* with his deciding to *A* at a time at which it is up to him what he decides. Alternatively, one can claim that an agent's deciding to *A* at a time at which it is up to him what he decides is necessary but not sufficient for his freely deciding to *A*.

Consider the following case. A master manipulator implants in Ken an irresistible desire to kill Larry within the next few minutes but leaves it to Ken to decide on the murder weapon. Ken has two options, an AMT hardballer (pistol 1) and an AMT longslide (pistol 2). Ken decides to kill Larry with the former weapon and does so. It was open to him to decide instead to kill Larry with the other gun. Various theoretical options are available, including the following.

- A. It was up to Ken whether he decided to kill Larry with pistol 1 or decided to kill him with pistol 2, and his deciding to kill Larry with pistol 1 was a free action.
- B. It was up to Ken whether he decided to kill Larry with pistol 1 or decided to kill him with pistol 2, and his deciding to kill Larry with pistol 1 was not a free action.
- C. It was not up to Ken whether he decided to kill Larry with pistol 1 or decided to kill him with pistol 2, and his deciding to kill Larry with pistol 1 was not a free action.

It is easy to imagine arguments for each of these options. I find B and C much more plausible than A; but rather than argue about this sort of thing, I have built into *UTn* agential properties that deflect cases of heavy-duty manipulation, massive deception, and the like.

6.7. Conclusion

I argued that the argument I examined from Clarke 2004 does not undermine the thesis that even if the difference at *t* in what an agent decides in a

pair of worlds with the same past up to t and the same laws of nature is just a matter of luck, the agent may directly freely decide what he decides. More colloquially put, the thesis is that the luck involved is compatible with making a directly free decision. I argued as well that the idea that this luck also is compatible with its being up to the agent what he decides merits consideration. Its being up to us what we decide may be interestingly different from its being up to us what we do when the options are overt actions (or mental actions of certain kinds: for example, silently reciting poems to oneself). The latter may require promise-level abilities (P -abilities). But its being up to me whether I decide to A or decide to B seems not to require that I am P -able to decide to A and P -able to decide to B . Promise-level ability is out of place when the topic is what an agent is able to decide to do.

So, is what you decide ever up to you? More specifically, is this ever up to you in a sense of that expression that requires your being able to make your decision directly freely? And even more specifically, is this ever up to you in such a sense even if you lack agent-causal powers? Maybe so. One thing it depends on is whether you sometimes are able to make directly free decisions. I say that the case for the claim that you are able to do this is stronger than the case for the claim that you are not, but my arguments for that assertion are elsewhere (Mele 2006).²⁰

Notes

1. For doubts about the possibility of agent causation, see Clarke 2003, chap. 10. For an argument against the existence of agent-causal powers, see Pereboom 2001, chap. 2. The thesis that agent causation does not solve the problems at issue in this chapter is defended in Mele 2006, chap. 3 (also see van Inwagen 2000).
2. Here I draw on a synopsis in Kearns and Mele 2014.
3. For a reply, see Mele 2006, pp. 53–56.
4. On complete control, see chapters 8 and 11.
5. Randolph Clarke argues that agent-causal powers are required for free will (at least, if incompatibilism is true), and in his judgment, relevant arguments collectively “incline the balance against the possibility of substance causation in general and agent causation in particular” (2003, p. 209). In conversation, Clarke said he had metaphysical possibility in mind.
6. Van Inwagen is open to there being different senses of “able” (2011, p. 482, n. 3).
7. I have noticed that some people are inclined to read “ S did not A ” as attributing to S an action of not A -ing. Hence, I insert “do” between “not” and “ A .”
8. Readers may be confused by the date of this publication. Clarke 2004 cites Mele 2006. Backward causation is one hypothesis. Late publication is another.

9. A comment on the quotation from Clarke is in order. Readers who believe that Al's voting for Gore in an earlier example is indirectly free and derives its freedom from the freedom of his pulling the Gore lever (see chap. 5, sec. 5.1) will reject the idea that a free action's not deriving its freedom from any *earlier* action is sufficient for its being directly free. Al's pulling the Gore lever and his voting for Gore begin at the same time, and there is no third, earlier action from which the freedom of his voting for Gore is thought to derive. (I should add that whatever view one takes on when actions end, the case can be presented in such a way that the actions at issue end at the same time. For example, for someone who holds that Al's voting for Gore does not end until the machine registers his vote, it can be a feature of the case that the lever's reaching a certain point along its path causes the vote to be registered just as the lever completes its motion.) Readers of the sort at issue here may hold that there are two different ways for an action to be indirectly free: by deriving its freedom from the freedom of an earlier action, and by deriving its freedom from the freedom of an action that begins at—or is performed at—the same time.
10. The inspiration for this story is Michael Bratman's well-known video games story (1987, p. 114).
11. Frankie Caruso (in a paper for a seminar I taught) suggested replacing *UT* with a principle that differs from it only in that "*I*-able" replaces "able." Caruso takes the view that Bart's raising his right index finger is not an intentional action and that Bart is not *I*-able at the time to raise it. The thought, of course, is that if Bart lacks that ability, my story does not falsify the modified version of *UT*. However, there is good reason to reject the claim that Bart's raising that finger is not an intentional action. Raising his right index finger is in Bart's basic action repertoire, he was trying to raise it, and he succeeded—in an utterly normal, nondeviant way—in raising it. Readers who find Caruso's contention initially plausible may be failing to distinguish the following two claims: (1) Bart's raising his right index finger is an intentional action; (2) It was intentional on Bart's part that he raised his right index finger rather than his left index finger.
12. Someone might erroneously claim that Bart's raising his right index finger derives its freedom from the freedom of another action—namely, his trying to raise it. Bart tries to raise the finger, and that attempt is successful. His trying to raise it is not an "earlier action" than his raising it. It is not as though he tries to raise it and then, after he tries, raises it. Nor is his trying to raise the finger a distinct action from his raising it that occurs simultaneously with his raising it. His raising the finger *is* his successful attempt to raise it (Adams and Mele 1992).
13. I am not claiming that Clarke's argument was intended to show this. Again, it is possible that he meant his argument to apply only to directly free decisions.
14. *UT* may be modified by inserting "at *t*" between "able" and "to." This change would raise an interpretive question. Is the modified version of *UT* only about cases in which an agent is able throughout *t* to *A* at *t* and able throughout *t* to *B* at *t*? Or is it also about cases in which at an early part of *t* (but not throughout *t*) the agent

is able to *A* at *t* and able to *B* at *t*? If the former, my story about Bart is not a counterexample to the modified version of *UT*: when his attempts begin, he becomes unable to raise his left index finger; he is unable to raise it for the remainder of *t*. But the scope of this modified version of *UT* is very restricted, and one would do well to set it aside and turn to the next item of business (*UTD*). I may be able now to jump to the east straightaway and able now to jump instead to the west straightaway; but once my feet leave the ground moving east, I am no longer able at that time to jump to the west just then.

15. *UI* is based on *P* above.
16. For the record, in Mele 2006, chap. 3, I argue against a claim about agent causation that is closely related to the one mentioned in the paragraph to which this note is appended; and in Mele 2013d, I develop a challenge for libertarians that avoids any mention of luck but is very similar to the problem about luck presented here (Mele 2006, chap. 3).
17. For evidence that a majority of lay folk use “up to [an agent]” in a way consistent with the idea that determinism does not preclude its sometimes being up to agents what they do, see Nahmias, Coates, and Kvaran 2007, p. 227.
18. Some may claim that libertarians cannot settle for this because the condition secures no more control than compatibilists can secure. For a critique of some control-based arguments for the thesis that event-causal libertarianism is an uninhabitable halfway house between compatibilism and agent-causal libertarianism, see chapter 7. Also see chapter 9, section 9.2.
19. Discussion with Stephen Kearns motivated this section.
20. For comments on a draft of Mele 2013b, on which much of this chapter is based, I am grateful to Randy Clarke and Stephen Kearns.

Arbitrary Decisions and the Problem of Present Luck

DURING HIS DISCUSSION of an argument about luck and free will that has been attributed to me, Bernard Berofsky writes: “What, according to Mele, would it be like *not* to be a recipient of luck (to have control in the sense required for free will)?” (2012, p. 67). The question is motivated by a mistake that I correct before I answer it. My focus in this chapter is on a positive element of a typical libertarian view: namely, the thesis (*LFT*) that there are indeterministic agents who sometimes act freely when their actions are not deterministically caused by proximal causes.¹ (The negative side of libertarianism, of course, is the thesis that free will is incompatible with determinism.) *LFT* is a target of what I call “the problem of present luck” (Mele 2006, p. 66) or, more fully, “the problem of present indeterministic luck” (p. 201). I sketched the problem in chapter 6 (section 6.2).

As I reported in chapter 1, my interest in free action is in what I called *moral-responsibility-level free action*—“roughly, free action of such a kind that if all the freedom-independent conditions for moral responsibility for a particular action were satisfied without that sufficing for the agent’s being morally responsible for it, the addition of the action’s being free to this set of conditions would entail that he is morally responsible for it” (Mele 2006, p. 17). *LFT* should be interpreted accordingly. *LFT* should also be distinguished from the related libertarian thesis (*LF*) that there are indeterministic agents who sometimes act freely. If there can be local deterministic causal connections in indeterministic universes, some actions performed by indeterministic agents may be deterministically caused by their proximal causes. If only such actions of indeterministic agents can be free, *LFT* is false.

7.1. Arbitrary Decisions

The question I quoted from Berofsky 2012 assumes that being “a recipient of luck” is incompatible with having “control in the sense required for free will.” But that is not an assumption I endorse. In fact, I reject it. And, despite some claims to the contrary, I have never offered a luck-based argument—or any argument at all—for the falsity of libertarianism. Instead, I have articulated a problem about luck for typical libertarians (2005; 2006, chap. 3), encouraged readers to find a solution (2005, 2006), and developed a solution of my own (2006, chap. 5). Regarding “the problem of present luck” (2006, p. 66), I wrote that “my aim in developing this chapter’s central problem for agent causationists and other conventional libertarians is to present it sufficiently forcefully to motivate them to work out solutions to it—proposed solutions that I and others can then assess” (2006, p. 70; see Mele 2005, p. 414). One plank in my proposed solution is the thesis that even if the difference between what an agent does at t in one possible world and what he does at t in another possible world with the same past up to t and the same laws of nature is just a matter of luck, the agent may perform a directly free action at t in both worlds (Mele 2006, chap. 5). I return to this thesis—*LDF* from chapter 6—in section 7.5.

In chapter 6, I illustrated the problem of present luck with a story (from Mele 2006) about a man named Bob who decided at noon to cheat. Given the details of Bob’s story, I asked, how can Bob have enough control over whether he decides to cheat or does something else instead at noon for his decision to be directly free and for him to be directly morally responsible for it? As I observe elsewhere (Mele 2013d, pp. 241–42), I regard my central question about stories of this kind as an analogue of a request for a theodicy in response to the problem of evil—an explanation of why a perfect God would allow all the pain and suffering that exists in the world. If I had wanted to prompt an analogue of a defense against an argument from evil for the nonexistence of God, I would have offered an argument from luck for the falsity of libertarianism and encouraged rebuttals. Rebuttals of such an argument might not have included answers to my question.

In one kind of scenario that Berofsky mentions in his discussion of the problem of present luck, “reason plays no role,” as in “the selection of a particular card in response to the magician’s request to pick a card” (2012, p. 65). He says that “although the choice of a card is arbitrary . . . the spectator has control over the selection” (p. 65), and he contends that “there is no reason to deny that the agent is choosing freely” (p. 66). Reflection on a case of this kind will prove useful.

I was once a participant in a neuroscience experiment of a well-known kind that is claimed to have a bearing on free will. My task was to flex my right wrist whenever I wished (on many occasions over the course of the experiment) while watching a very fast clock and then to report where the hand was on the clock when it seemed to me that I decided to flex. I had no reason to prefer any moment to begin flexing over any nearby moment. So I was in the sort of situation Berofsky has in mind. The magician asks someone from the audience to pick a card—any card—and I was asked to pick a moment to begin flexing.

The strategy I employed was to say “now!” silently to myself and treat that silent speech act as an expression of a proximal decision to flex for the purpose of having something to report. I have no idea why I said “now!” exactly when I did rather than a few—or even many—milliseconds earlier or later. Consider a particular “now!”-saying of mine. Call the time at which it happened t . If it was at no time determined that I would say “now!” at t , then there are other possible worlds with the same laws and the same past up to t in which I do not say “now!” at t . In some such worlds, I continue watching the clock and flex a bit later. If this is how things were, I satisfied a necessary condition for directly free action according to typical libertarians. But, keeping in mind that by free action I mean moral-responsibility-level free action, did I freely say “now!”? When I said “now!” I exercised some control, but was it moral-responsibility-level control? (I am not asking whether I was morally responsible for that silent speech act; see chapter 1, note 3.)

These questions are interesting, but I will not try to answer them here. I raise them as background for a pair of observations about potential consequences of focusing on cases of arbitrary picking when considering the problem of present luck. (Berofsky does not focus exclusively on such cases. The points I am about to make are not criticisms of his work.) One potential consequence is the activation of a conception of free action that falls short of moral-responsibility-level free action. People—philosophers and others—may conceive of free will somewhat differently on different occasions, depending on what kind of scenario they are considering; and, in the same person, a conception of free will activated by consideration of a case of one kind may be less demanding than a conception of free will activated by consideration of a case of another kind.

Another potential consequence is that one may ignore a route to motivating the problem of present luck that has some intuitive clout. My preferred route goes through moral responsibility. Recall the “fuller version” of Bob’s story (chapter 6, section 6.2), in which he does his best to resist temptation.

In a discussion of the problem it raises, I would do the following two things. First, I would motivate the idea that Bob does not deserve to be blamed for his bad decision. (This is not to say that I endorse this idea. Again, I believe that the problem of present luck has a solution.) Then I would argue that, given the details of the case, if he made that decision freely, he would deserve to be blamed for it (see Mele 2006, pp. 63–65 on a related story). But in typical cases in which the best an agent can do is to pick something or other arbitrarily, the issue of deserved blame does not even arise.

Pretend that my brain works in the following way when I make arbitrary decisions on which nothing significant hangs in situations in which I am committed to *A*-ing, the time at which I *A* is left up to me, and I have agreed not to *A* unless I first make a proximal decision to *A*.² When I am informed that it is time to get ready to make a decision, a little neural “decision wheel” is activated in my head, and a tiny neural ball starts bouncing along it. The wheel has tiny slots on it. They are all the same size, like the slots on a roulette wheel. When the wheel is activated, there is an objective probability that I will make my decision at moment *m*₁, an objective probability that I will make it at *m*₂, and so on; and the different probabilities are correlated with the number of slots assigned to a decision made at a particular moment. Finally, my making the decision to *A* that I make—a token event—is the ball’s landing in a certain slot. (Or, if you prefer, the neural realizer of my making that token decision is the ball’s landing there.)

As a subject in the experiment I described, I was arbitrarily selecting a moment to begin flexing. The decision wheel seems quite fine for my decision task. If I were to discover that I work this way as a decision-maker in cases of this particular kind, I would not worry that there are significantly better ways to make these decisions.

Pretend now that we discover that, *whenever* we make a decision, some decision wheel is at work. In stories like Bob’s, where different action-types are options, objective probabilities of the various possible outcomes are reflected in the number of slots assigned to an outcome. In Bob’s case, a collection of slots represents his deciding at noon to toss the coin straightaway, and collections of other slots represent various other outcomes. The wheel is activated by Bob’s desire to decide the matter, the tiny ball starts bouncing along it, and what decision he makes is a matter of where the ball lands.

In this scenario, it looks like there *is* something to worry about. If this is how Bob works, does he have sufficient control over whether he makes a bad decision or does something else instead to make his decision freely and to be morally responsible for the decision he makes? Is what Bob decides too

much a matter of chance for him to make his decision freely? I do not answer these questions in this chapter (but see Mele 2006, chap. 5). My aim in this section has been to support a pair of claims about potential consequences of focusing on cases of arbitrary picking when considering the problem of present luck: namely, that such a focus might activate a conception of free action that falls short of moral-responsibility-level free action and that it might lead one to ignore a route to motivating the problem of present luck that has some intuitive force.

7.2. A Response to Berofsky's Question for Me

I return to the question quoted at the outset of this chapter: "What, according to Mele, would it be like *not* to be a recipient of luck (to have control in the sense required for free will)?" (Berofsky 2012, p. 67). As I mentioned, I reject the assumption that being "a recipient of luck" is incompatible with having "control in the sense required for free will." So I view Berofsky's question as a pair of questions.

- Q1. "What, according to Mele, would it be like *not* to be a recipient of luck?"
 Q2. "What, according to Mele, would it be like . . . to have control in the sense required for free will?"

I respond to each in turn.

I have never given a definition of "luck." But, in *Free Will and Luck* (Mele 2006), I did say the following about the luck that primarily concerned me there:

What, I will be asked, is luck? Well, if the question why an agent exercised his agent-causal power at *t* in deciding to *A* rather than exercising it at *t* in any of the alternative ways he does in other possible worlds with the same past and laws of nature is, in principle, unanswerable—unanswerable because there is no fact or truth to be reported in a correct answer, not because of any limitations in those to whom the question is asked or in their audience—and his exercising it at *t* in so deciding has an effect on how his life goes, I count that as luck for the agent—good luck or bad, depending on the goodness or badness of the effect the particular exercise of agent-causal power has. If "luck" is not the best short label for this sort of thing, I am open to correction. Whatever it is called, agent causationists should try to persuade people

who have the worry I have described that this worry should not stand in the way of their accepting agent-causal libertarianism. (2006, p. 70)

Agent causation is in the forefront here because at this point in the book I had been arguing that something like the “luck” or “chance” objection that some fans of agent causation raise against event-causal libertarianism applies also to agent-causal libertarianism. In any case, in response to Q1, I can say that a sufficient condition for not being “a recipient of” the luck that primarily concerned me—indeterministic luck at the time of decision or at the time of a pertinent exercise of agent-causal power—is that the decision or exercise was deterministically caused by its proximal causes. This is not, of course, a general answer to Q1; the answer is strictly about the luck that primarily concerned me in Mele 2006.³ It should be noted that on the assumption that there can be local deterministic causal connections in indeterministic worlds, the answer I just offered is distinguishable from the following answer: the decision or exercise occurred in a deterministic world.

I turn to Q2. In Mele 2006, I developed both a compatibilist view and a libertarian view while remaining agnostic about compatibilism (both about free will and about moral responsibility). In this connection, in the book’s final chapter, I offered some different sufficient conditions for an agent’s freely *A*-ing. One was a compatibilist condition featuring an “ideal agent”:

1a. An agent *A*-s freely if he nondeviantly *A*-s on the basis of a rationally formed deliberative judgment that it would be best to *A*, he is an ideally self-controlled and mentally healthy person who regularly exercises his powers of self-control, he has no compelled or coercively produced attitudes, his beliefs are conducive to informed deliberation about all matters that concern him, and he is a reliable deliberator. (2006, p. 200)

I observed that if 1a is true, “the door is wide open to realistic versions” of it, such as the following one:

1b. An agent *A*-s freely, if he nondeviantly *A*-s on the basis of a rationally formed deliberative judgment that it would be best to *A*, has no compelled or coercively produced attitudes that influence his deliberative judgment, is well-informed on the topic of his deliberation, and is mentally healthy. (2006, p. 200)

I also offered libertarians two different sufficient conditions for freely deciding to *A*, one building on 1a and the other on 1b:

2a. An agent freely decides to *A* if he nondeviantly decides to *A* on the basis of a rationally formed deliberative judgment that it would be best to *A*, the proximate causes of his decision do not deterministically cause it, and he satisfies the conditions stated in 1a as fully as that is possible given that the proximal causes of his decision do not deterministically cause it. (2006, p. 201)

2b. An agent freely decides to *A* if he nondeviantly decides to *A* on the basis of a rationally formed deliberative judgment that it would be best to *A*, the proximate causes of his decision do not deterministically cause it, he has no compelled or coercively produced attitudes that influence his deliberative judgment, he is well-informed on the topic of his deliberation, and he is mentally healthy. (2006, p. 201)

All four of these conditions are associated with an answer to Q2. I believe that if compatibilism is true, then any agent who satisfies condition 1a or 1b has all the control required for acting freely. I also believe that any agent who satisfies condition 2a or 2b has all the control required for deciding freely.⁴ Given that I believe what I do about the two libertarian conditions, readers are entitled to infer that I believe that the problem of present luck—including indeterministic luck at the time of *decision*—is soluble. And, as I mentioned, I offered a solution to it in Mele 2006. The solution is not a response to an argument, because I offered no argument for the falsity of libertarianism. It is a response to the problem sketched in chapter 6 (section 6.2).

Readers should bear in mind that the numbered conditions are proposed *sufficient* conditions. It should not be inferred, for example, that I believe that acting in accordance with what one judges best is a *necessary* condition for acting freely. In fact, I hold that agents can decide freely—and, more generally, act freely—contrary to what they judge best (Mele 2006, pp. 118–29).

7.3. Event-Causal Libertarianism and Control

One response to conditions 2a and 2b above is that no real libertarian would be satisfied with them. Derk Pereboom claims (*DP*) that event-causal libertarianism fails because it “does not provide agents with any more control than compatibilism does” (2001, p. 56). And Randolph Clarke argues (*RC1*) that

“the active control that is exercised on [an event-causal libertarian] view is just the same as that exercised on an event-causal compatibilist account,” adding (*RC2*) that the “view fails to secure the agent’s exercise of any further positive powers to causally influence which of the alternative courses of events that are open will become actual” (2003, p. 220).

We have in *RC1* a partial basis for a control-featuring argument against event-causal libertarianism. I call it the *same-control argument*:

SL. Having free will depends on having a kind of active control that cannot be had in deterministic worlds and therefore cannot be captured by compatibilist accounts of free will.

RC1. “The active control that is exercised on [an event-causal libertarian] view is just the same as that exercised on an event-causal compatibilist account.”

So *S3*. Event-causal libertarianism is false.

One might take one’s lead in trying to ascertain whether *RC1* is true from what John Fischer refers to as the distinction between “guidance control” and “regulative control” (1994, pp. 132–35). Guidance control can be exercised in deterministic worlds. An example is the control we normally exercise over how the cars we are driving move, if our world is deterministic. But regulative control, by definition, cannot be exercised in any deterministic world. When one *A*-s at *t*, one exercises regulative control over one’s *A*-ing only if there is another possible world with the same laws of nature and the same past up to *t* in which one does not *A* at *t*. Regulative control is not “just the same as” guidance control. And the active regulative control that is exercised on an event-causal libertarian view is not “just the same as” the active guidance control that is “exercised on an event-causal compatibilist account.” The former, by definition, is incompatible with determinism and the latter, by definition, is compatible with determinism. These facts preclude their being “just the same.” It can be said that event-causal regulative control is “just the same as” event-causal guidance control in a certain obvious respect—both are event-causal. But we need an argument to show us why we should believe that some true control requirement for free action demands something that no event-causal view of control can provide.

Even though it is false that regulative control is “just the same” as guidance control, we still have the issue about “more control” that Pereboom raises (2001, p. 56). In discussions of comparative control in the free will literature, *direct control* is a prominent notion. Clarke writes: “Direct active control is

exercised in acting, not before" (2003, p. 166). O'Connor reports that "exerting active power is intrinsically a direct exercise of control over one's own behavior" (2000, p. 61). And Robert Kane claims that agents exercise direct control over some of their choices (1996, p. 144). In these cases, Kane says, the agent's exercise of control is not "antecedent" to the choice; rather, it occurs "then and there," when and where the choice is made.

The following argument features direct control. I dub it the *more-control argument*.

*M*₁. Necessarily, if an agent's world is deterministic, then even if he has as much control as an agent can possibly have in a deterministic world, he lacks free will.

*M*₂. Necessarily, an agent with no agent-causal powers who has as much direct indeterministic control as can be had in the absence of agent-causal powers does not have a greater amount of control than an agent who has as much control as can be had in a deterministic world.

So *M*₃. Necessarily, even an agent with as much direct indeterministic control as can be had in the absence of agent-causal powers lacks free will if he has no agent-causal powers.

If this argument is to be valid, it needs another premise. To see why, suppose that the following proposition is true:

P. An agent with direct indeterministic control and no agent-causal powers can have free will even if he does not have a greater amount of control than an agent who has as much control as can be had in a deterministic world.

There is no explicit contradiction in the conjunction of *M*₁, *M*₂, and *P*; but *P* entails that *M*₃ is false.

As far as I know, no one who argues that event-causal libertarianism is false on the basis of considerations of the sort featured in premises *M*₁ and *M*₂ has told us how to measure amounts of control or how to weigh deterministic and indeterministic control on the same scale.⁵ But I am willing to grant premise *M*₂ of the more-control argument for the sake of argument. If the premise is true, the agents at issue might have the same amount of control or it may be that direct indeterministic control and any kind of deterministic control are incommensurable. In any case, in the absence of an argument against *P*, the more-control argument is at best incomplete.

Consider the claim that if two cars do not differ in horse power, the top speed of either cannot be greater than that of the other. This claim is false. Other features of cars are relevant to how fast they can move. Might amounts of control be like that in the sphere of free will? Someone might claim that if all relevant features of two agents that are not control features are equal, then if the agents do not differ in the amount of control they exercise at a time, either both act freely at that time or neither does. An argument for this claim may prove illuminating.

Premise *Mr* of the more-control argument—the incompatibilist premise—is relevant in this connection. A philosopher who assents to it may say that agents in deterministic worlds do not exercise enough control when they act to act freely, but the same philosopher may add that these agents do not exercise enough control because they do not have the right kind of control. What is the right kind? According to some philosophers, having the right kind of control requires having agent-causal powers—powers that they themselves are inclined to regard as impossible (Clarke 2003, p. 209) and powers the existence of which they say we have no evidence for (Clarke 2003, pp. 206–7) or weighty evidence against (Pereboom 2001, chap. 3).⁶ According to others, the right kind of control is a species of direct indeterministic control that is unsupplemented by any agent-causal powers (Kane 1996). An event-causal libertarian may claim that an agent can exercise enough of this kind of control to act freely even if the amount of direct control he exercises does not surpass the greatest amount of control open to agents in deterministic worlds. (Instructions about how to weigh amounts of deterministic and indeterministic control on the same scale might prove useful for those who wish to assess this claim.)

Conceptual sufficiency and contingent sufficiency are very different, and both are sometimes talked about in terms of what is *enough* for something. I recently learned that three cases of beer are not enough for a party attended by certain friends of mine and that five cases are enough. This is a contingent matter. I also learned long ago that for something to be a line, it is enough that it be a curve. This is a matter of conceptual sufficiency.

In light of the simple distinction just mentioned, it is easy to see an ambiguity in the claim that Andy did not exercise *enough* control to act freely. On one reading, the claim is that the control he exercised is not part of something *conceptually sufficient* for an action's being free. If Andy satisfied all necessary conditions for having *A*-ed freely that are independent of control, he might have failed to act freely because he did not exercise a certain *kind* of control. On the reading currently under consideration, one may try to defend the

claim that Ann, unlike Andy, did exercise *enough* control to have acted freely without saying anything about relative *amounts* of control. What matters, one may think, is that Ann exercised a species of direct indeterministic control whereas Andy exercised only deterministic control.

On another reading of the claim about Andy, what is being asserted is that the amount of control he exercised falls short of some amount required for free action. Similarly, on a reading of this kind, the claim about Ann is that she exercised a greater amount of control than Andy did and an amount great enough for her to have acted freely. Someone who prefers these readings of the claims about Ann and Andy should tell us how to measure amounts of control and how to weigh deterministic and indeterministic control on the same scale.

We all learned in school that being a curve is sufficient for being a line. This sufficiency is not a matter of an amount of anything; it is simply a matter of definition. Might it be, similarly, that a particular exercise of direct indeterministic control is part of something conceptually sufficient for a particular action's being a free action although even a maximal exercise of deterministic control cannot play this role and the two exercises differ in kind but not in amount? Perhaps, in the category of amount, they are incommensurable.

One might reply as follows: Other things being equal, if the former exercise of control was sufficient to play the role at issue but the latter was not, then the former must have involved more control than the latter. But this is simply to repeat a thought that I have been challenging. Would someone who has this thought feel compelled to explain the following comparative fact about sufficiency in terms of different amounts of something? Although being a curve is sufficient for being a line, being a tomato is not.

Consider the following argument. I dub it the *lame-control argument*.

L1. No exercise of event-causal control, no matter how impressive, is an exercise of freedom-level control in any deterministic universe.

So *L2.* No exercise of event-causal control is an exercise of freedom-level control in any indeterministic universe.

The obvious invalidity of this argument may explain why I have never seen it. A premise like the following would fill the gap:

Lb. *L1* is true because no possible exercise of event-causal control is an exercise of freedom-level control.

(*Lb* would fill the gap in an odd way, given that *Lb* would be doing all the work.) If either the same-control argument or the more-control argument were successful, it would support *Lb*. But, for all these arguments show, if *L1* is true, it is true because determinism precludes the exercise of freedom-level control and not because of any general, determinism-independent fact about event-causal control. Clarke has objected to event-causal libertarianism on the grounds that it adds no “positive” power of control to what compatibilists can offer but simply places compatibilist control in an indeterministic setting (2000, p. 35). In Mele 2006 (p. 14), I observe that, given that placing event-causal control in an indeterministic setting was my explicit strategy in Mele 1995a for generating a libertarian position on sufficient conditions for free action, I do not see this as an objection. I have the same attitude now toward the lack of an additional control power—a non-event-causal power. No one has built a successful bridge from *L1* to *L2*. And no successful argument for *L2* exists (for more on this see chapter 8).

7.4. Agent Causation and Control

Does regulative control “provide agents with any more control” than guidance control does? Instructions about how to measure amounts of control and how to weigh guidance and regulative control (or deterministic and indeterministic control) on the same scale would seem to be an essential part of a persuasive argument for an answer. Yet, as I have observed, philosophers who have answered *no* to the question have provided no such instructions. Some such philosophers agree, however, about where to look for more control. Pereboom holds that agent causation, if it were to exist, would provide the “enhanced control” for which he calls (2001, p. 55); and Clarke contends that “the requirement of agent causation . . . provides for the agent’s exercising when she acts, in addition to the active control secured by an event-causal view, a further power to causally influence which of the open alternatives will be made actual” (2003, pp. 220–21).

Both Pereboom and Clarke are skeptical about agent causation—and rightly so, in my view (Mele 2006, chap. 3). Pereboom argues that although agent causation is possible, it is extremely unlikely that anyone has agent-causal power (2001, chap. 3). This is part of his argument for the thesis that no one has free will. In Clarke’s judgment, relevant arguments collectively “incline the balance against the possibility of substance causation in general and agent causation in particular” (2003, p. 209), and he contends that there is no evidence for the existence of agent causation (pp. 206–7). In a recent book,

Pereboom abandons his claim that agent causation is possible—conceptually and metaphysically—in favor of agnosticism about this (2014, p. 58).

One route to skepticism about free will features the combination of incompatibilism and the contention that free will depends on agent causation. Some arguments for the contention about agent causation include such unsubstantiated claims as that “the active control that is exercised on [an event-causal libertarian] view is just the same as that exercised on an event-causal compatibilist account” (Clarke 2003, p. 220) and that event-causal libertarianism “does not provide agents with any more control than compatibilism does” (Pereboom 2001, p. 56). In Mele 2006, I suggest, in effect, that libertarians would do better to think in terms of *kinds* of control than in terms of *amounts* of control. (To forestall confusion, I point out that I do not place any special weight on the word “kinds” here. I need a term to contrast with “amounts” in the sphere of control, and “kinds” seems to be a reasonable choice.) In the absence of instructions about how to measure amounts of control and how to weigh indeterministic control and deterministic control on the same scale, this seems wise. One question event-causal libertarians face in this connection is whether there is a sufficient condition for free will (or free action) that features indeterministic control and makes no appeal to agent causation. Another is whether the sufficient condition they offer is superior to the most promising sufficient condition for free will (or free action) that is compatible with determinism.

Before I take up these questions, a comment is in order about the spirit in which I do so. Early in *Nature’s Challenge to Free Will*, Berofsky writes “I want to understand, even empathize with my opponents. I may not change minds; but if I am successful, we will understand better what our competing outlooks really are and what the nature of a resolution to our dispute looks like” (2012, p. 1). I have a similar attitude. In fact, an unusual feature of my position is that it keeps both compatibilism and event-causal libertarianism in the running (Mele 1995a, 2006). I empathize with proponents of both of these two opposing views, and I attempt to solve theoretical problems for both groups. I have no interest in defending the negative side of event-causal libertarianism—that is, incompatibilism. But I find the positive side very interesting—especially *LFT* (identified in this chapter’s introduction and repeated shortly).

Berofsky is a compatibilist, and I am officially agnostic about compatibilism. He and I may each find some compatibilist sufficient condition for free will and free action very attractive (perhaps the same condition and perhaps not). And he may believe, as I do, that adding a certain element of 2a and 2b to our proposed compatibilist sufficient condition for free decisions—namely,

that the proximate causes of the decision do not deterministically cause it—would not necessarily result in a condition that falls short of being sufficient for free decisions (because, for example, satisfying it precludes deciding freely). But some incompatibilists may claim that Berofsky and I set an unacceptably low bar for free will and free decisions and that any right-thinking incompatibilist would see that adding just the element mentioned to a compatibilist condition that he or I find attractive (for example, adding it to 1a or 1b) would not carry one over the true bar for free decisions. I have not yet seen a convincing argument for this claim. But the absence of a convincing argument for it does not warrant the belief that the claim can safely be ignored.

Where the bar should be set for free will is the subject of much controversy, of course, and a great many competing options have been offered. If there were a prize for the most outlandish option from the recent scientific literature on free will, I might nominate the following passage from an article entitled “Free Will” by neuroscientist P. Read Montague. I quoted it in chapter 5, but it is so wonderfully over the top that I quote it again here:

Free will is the idea that we make choices and have thoughts independent of anything remotely resembling a physical process. Free will is the close cousin to the idea of the soul—the concept that “you,” your thoughts and feelings, derive from an entity that is separate and distinct from the physical mechanisms that make up your body. From this perspective, your choices are not caused by physical events, but instead emerge wholly formed from somewhere indescribable and outside the purview of physical descriptions. This implies that free will cannot have evolved by natural selection, as that would place it directly in a stream of causally connected events. (2008, p. 584)

Obviously, no one who sets the bar for free will where Montague does will deem any of 1a through 2b true. And the closer one’s bar is to this outlandish one, the less attractive these conditions will seem to one. The same goes for the solution to the problem of present luck that I offered libertarians. It is a solution for libertarians who believe that all actions are caused and therefore would be regarded as a nonstarter by Montague. (The solution is developed in a lengthy stretch of chapter 5 of Mele 2006, and I reply to some objections in Mele 2013d, pp. 253–54. I have more to say about it in chapter 10.) Between Montague’s bar and bars that compatibilists find attractive there are lots of options. Some of the proposed bars feature both agent causation and indeterminism. I mentioned Pereboom’s agnosticism about the metaphysical

and conceptual impossibility of agent causation (also see my quotation from Clarke 2003, p. 209). If agent causation is impossible, then any bar for free will that includes agent causation is impossibly high. The distance between an impossibly high bar and something like 1b, a compatibilist proposition, is incalculable. And the same is true when something like 2b, an event-causal libertarian proposition, is substituted for 1b. This may help to explain why some philosophers who contend that free will depends on agent causation are unimpressed by differences between event-causal guidance control and event-causal regulative control. Because both kinds of control are incalculably distant from any impossible species of control, they may appear to be so close to each other that any difference between them cannot have an interesting bearing on free will.

Not everyone who contends that free will depends on agent causation voices serious doubts about the conceptual or metaphysical possibility of agent causation, as I have mentioned. In any case, the issue about perspective that I raised is just one issue relevant to a discussion of bar setting. Another is the quality of the arguments for setting the bar for free action where one does.

Consider the following claim: (*C*) If compatibilism about free action is false, then event-causal libertarianism is false. How might it be defended? One premise in one kind of argument for *C*—a premise that would also require a defense—is the claim (*CI*) that the control available to agents in any deterministic world cannot satisfy some control requirement for free action. I have already mentioned two options for a second premise:

(*RCI*) “the active control that is exercised on [an event-causal libertarian] view is just the same as that exercised on an event-causal compatibilist account” (Clarke 2003, p. 220).

(*DP*) Event-causal libertarianism “does not provide agents with any more control than compatibilism does” (Pereboom 2001, p. 56).

In section 7.3, I explained why *RCI* is unacceptable, and I rebutted an argument against event-causal libertarianism that appeals to *DP*.

Return to *C*. Obviously, a compatibilist may agree that *C* is true and reject the antecedent. If compatibilism is true, then the negative side of event-causal libertarianism—the assertion of incompatibilism—is false. But for all anyone has demonstrated, the positive side of event-causal libertarianism—the thesis (*LF*) that there are indeterministic agents who sometimes act freely—may be true. Typical present-day compatibilists do not claim that free will requires determinism. So it seems that typical present-day compatibilists who believe

that we are indeterministic agents would accept *LF*. Such compatibilists have a stake in resisting threats to *LF*, just as event-causal libertarians do. (Threats to the idea that there are indeterministic agents need not worry compatibilists, of course. But *LF* asserts more than that.)

The problem of present luck targets a more specific thesis than *LF*, the thesis (*LFT*) that there are indeterministic agents who sometimes act freely when their actions are not deterministically caused by proximal causes. Any compatibilists who believe that *LFT* is true have a reason to tackle the problem of present luck. The same is true of compatibilists who have advanced a sufficient condition for deciding freely that is satisfied in some cases in which a decision is not deterministically caused by its proximal causes.⁷ Of course, a compatibilist who believes that the problem is insoluble can modify the sufficient condition at issue accordingly.

I identified a pair of questions event-causal libertarians face in a context in which it is claimed that agent causation is essential to a solution to the problem of present luck. First, is there a sufficient condition for free will (or free action) that features regulative (or indeterministic) control and makes no appeal to agent causation? Second, if so, is this sufficient condition superior to the most promising sufficient condition for free will (or free action) that is compatible with determinism? In Mele 2006, I argue for an affirmative answer to the first question, but without claiming that the condition is satisfied by actual human agents. Given my agnosticism about compatibilism, it is unsurprising that I take no stand on the second question. But notice that 2b above is a *different* condition from 1b, and a major difference is pertinent to the dispute between compatibilists and incompatibilists. 1b, but not 2b, can be satisfied in a deterministic world. It certainly is open to a libertarian to accept 2b while rejecting 1b and to argue that the former is superior to the latter.

Here is 2b again: “An agent freely decides to *A* if he nondeviantly decides to *A* on the basis of a rationally formed deliberative judgment that it would be best to *A*, the proximate causes of his decision do not deterministically cause it, he has no compelled or coercively produced attitudes that influence his deliberative judgment, he is well-informed on the topic of his deliberation, and he is mentally healthy” (Mele 2006, p. 201). I know of no convincing argument that 2b is false. My proposed solution to the problem of present luck is, among other things, a proposed solution to a worry about 2b. I continue to believe that “the following disjunction is more credible than the thesis that no human beings ever act freely and morally responsibly: either compatibilism is true and there are free and morally responsible human beings or compatibilism is false and there are free and morally responsible human beings” (Mele 2006, p. 206).

7.5. Parting Remarks

I mentioned that, in my view, even if the difference between what an agent does at t in one possible world and what he does at t in another possible world with the same past up to t and the same laws of nature is just a matter of luck, the agent may perform a directly free action at t in both worlds. This is the thesis I labeled *LDF* in chapter 6. Someone might ask how, given that I believe *LDF*, I can believe that there is a problem about luck for libertarians. The question puts the cart before the horse. Once the problem of present luck is made salient, one can get to work on finding a solution. Some philosophers look to agent causation for the solution, as I mentioned. And some such philosophers may be thinking that a decision's involving agent causation precludes the involvement of luck in the making of that decision even in cases in which, in another possible world with the same past and laws, the agent makes an alternative decision at that time. They may seek to solve the problem of present luck by showing that agent causation eliminates the luck at issue.

Meghan Griffith is one such philosopher. Central to her reply to the problem of present luck is the following claim, which she dubs RUL: "In the free will context, something is a matter of luck for A if and only if it happens to A " (2010, p. 46). She writes: "RUL cashes in on the distinction between doing and happening. What happens to someone is not something she does. In the free will context, that something happens to A is sufficient for its being a matter of luck, since it rules out her having done it. Since this is ruled out, her responsibility for it is undermined" (p. 46). In Griffith's view, decisions involving exercises of agent-causal power do not happen to the agents and therefore are not even partly a matter of luck.

In chapter 8, I argue that RUL is false. Furthermore, as I see it, present luck is ineliminable from scenarios of the kind at issue, and appeals to agent causation leave the problem of present luck unsolved. But that is not to say that there is no solution. I offer one in Mele 2006 (chap. 5), and *LDF* is part of it.⁸

Notes

1. Notice that *LFT* does not assert that actions have proximal causes. In principle, *LFT* can be endorsed by philosophers who hold that all free actions are uncaused, by agent-causal libertarians, and by event-causal libertarians.
2. For a related thought experiment, see Mele 2006, pp. 8–9.
3. It should not be inferred that I believe that no luck other than the luck that primarily concerned me poses an apparent threat to free will. For a problem about luck in deterministic scenarios that resembles in some respects what I call "the problem of

present luck” (Mele 2006, p. 66), see Levy 2011; Pérez de Calleja 2014; and Mele 2015b. For what I call the problem of “remote deterministic luck,” see Mele 2006, pp. 77–78.

4. As I reported in chapter 1, I conceive of free will as the ability to act freely. As I see it, any agent who acted freely at a time was able to act freely at that time and therefore had free will at that time.
5. For a brief comparative discussion of amounts of indirect control in a particular connection, see Mele 2006, pp. 62–63.
6. Not all philosophers who claim that free will depends on agent causation are skeptics about agent causation. See O’Connor 2000.
7. For replies by compatibilists to the problem of present luck, see Fischer 2012, chap. 6, and 2014; and Vargas 2012. For a critical assessment of these replies, see Kearns and Mele 2014.
8. Much of this chapter is based on Mele 2015c, a version of which I presented at Columbia University in March 2015. I am grateful to the audience for discussion. My critique of the more-control argument derives from Mele 2013c. I am grateful to Randy Clarke and Stephen Kearns for comments on a draft of the latter article.

Complete Control and Disappearing Agents

IN CHAPTER 7, I explained why two familiar control-based objections to event-causal libertarianism are unpersuasive. In the present chapter, I take up a third argument of this kind—Derk Pereboom’s “disappearing agent objection” (2014, p. 32). It runs as follows:

The disappearing agent objection: Consider a decision that occurs in a context in which the agent’s moral motivations favor that decision, and her prudential motivations favor her refraining from making it, and the strengths of these motivations are in equipoise. On an event-causal libertarian picture, the relevant causal conditions antecedent to the decision, i.e., the occurrence of certain agent-involving events, do not settle whether the decision will occur, but only render the occurrence of the decision about 50% probable. In fact, because no occurrence of antecedent events settles whether the decision will occur, and only antecedent events are causally relevant, *nothing* settles whether the decision will occur. Thus it can’t be that the agent or anything about the agent settles whether the decision will occur, and she therefore will lack the control required for basic desert moral responsibility for it. (2014, p. 32)

Pereboom adds: “The concern raised is that because event-causal libertarian agents will not have the power to settle whether the decision will occur, they cannot have the role in action that secures the control that this sort of moral responsibility demands” (p. 32).

Pereboom's disappearing agent objection is an important plank in his "case for free will skepticism" (2014, p. 71). It is intended to knock event-causal libertarianism out of contention. In this chapter, I explain why it fails to do that while also exploring a notion of complete control over whether one will decide to *A*.

8.1. Some Background

Pereboom's disappearing agent objection features a situation of a very specific kind. In it, the competing motivations of an agent who decides to *A*—"moral motivations" and competing "prudential motivations"—are "in equipoise" and "relevant causal conditions antecedent to the decision . . . render the occurrence of the decision about 50% probable" (2014, p. 32). Suppose, for the sake of argument, that the objection proves that an agent in this particular condition lacks "the control required for basic desert moral responsibility" for his decision to *A* because "nothing settles whether the decision will occur" (p. 32). What about agents with competing motivations that are not in equipoise? And what about agents whose competing motivations at the pertinent time do not include, for example, moral motivations? The objection at issue (which I quoted in full) makes no explicit claims about decisions these agents make. Yet, if Pereboom's disappearing agent objection is to falsify event-causal libertarianism, it must show that "event-causal libertarian agents" who make these decisions lack "the control required for basic desert moral responsibility" for them (p. 32).

Perhaps Pereboom would endorse the following claim: (*P_I*) If "event-causal libertarian agents" in the precise situation he describes lack "the control required for basic desert moral responsibility" for their decisions because nothing settles whether the decisions will occur (2014, p. 32), then any event-causal libertarian agent in a situation in which "relevant causal conditions antecedent to the decision . . . render the occurrence of the decision" less than 100% probable is in the same boat regarding the absence of settling and the control at issue.

P_I builds a bridge between Pereboom's narrowly focused disappearing agent objection (as quoted above) and the general conclusion about event-causal libertarianism for which he means to be arguing. Is *P_I* plausible? And, for that matter, is the much narrower claim specifically about agents whose competing motivations—moral versus prudential—are in equipoise plausible? That depends, among other things, on whether it is plausible that an agent's settling whether his decision to *A* will occur is required for his having

basic desert moral responsibility for the decision. I confess to lacking a good grip on what Pereboom means by “settles” in the expression “settles whether the decision will occur.”¹ This makes it difficult for me to assess *P1* and Pereboom’s disappearing agent objection itself.

As I observe elsewhere, “‘Decision’ has multiple referents. It refers (1) to the act of deciding; (2) to the immediate issue of the act, a decision *state*, a state of being decided upon something; and (3) to *what* we decide, as in ‘Her decision was to *A*’” (Mele 1992a, p. 158).² As I see it, “in deciding to *A*, one *settles* upon *A*-ing (or upon trying to *A*), and one enters a state—a decision state—of *being settled upon A*-ing (or upon trying to *A*)” (Mele 1992a, pp. 158–59). Presumably, a person’s settling on *A*-ing (in deciding to *A*) is different from a person’s settling whether a decision to *A* will occur, if what is meant by a decision to *A* is an act of deciding to *A*. I have a good grip on settling of the former kind, I believe (see chapter 2); but I cannot say the same about settling of the latter kind.

Suppose that Joe will decide at noon to skip his 2:00 class. What would it be for him to *settle* whether this decision will occur (as opposed to simply settling on skipping the class)? If Joe is, in Pereboom’s words, an “event-causal libertarian agent” (2014, p. 32), then, if Pereboom is right, he cannot settle whether this decision will occur. But what is this thing that he cannot do?

Pereboom does not answer these questions. But he does say the following: “If only events are causes and the context is indeterministic, the agent disappears when it needs to be settled whether the decision will occur, while the power of the agent to substance-cause decisions can have this settling role” (2014, p. 55). Pereboom is unsure whether the power just mentioned is metaphysically or even conceptually possible (p. 58). Even so, one may expect that by attending to what he says about this power and exercises of it, one can get a sense of what he means by an agent’s “settling” whether a decision will occur.

I will reproduce two relevant passages from Pereboom’s discussion of agent causation shortly. They appear in his discussion of a worry about luck that Ishtiyaque Haji (2004) and I (Mele 2005; 2006, chap. 3) raise for agent-causal libertarians. I set the stage for Pereboom’s discussion of that topic with some background on the worry. As I observed in chapter 6, some agent-causal libertarians contend that their event-causal cousins face an important problem about luck (or chance) in the case of decisions.³ In Mele 2006 (p. 54) I formulate the problem in terms of cross-world differences at the moment of decision. For example, if, in the actual world, Jim (who does not have agent-causal power) decides at *t* to keep working and in another possible world with the same past up to *t* and same laws of nature, Jim decides at *t* to take a break,

the difference in what Jim decides in the two worlds seems to be just a matter of luck, in which case his deciding at t to keep working seems to be partly a matter of luck. Philosophers have worried that luck in this connection is an obstacle to an agent's deciding freely and being morally responsible for his decision.

I have offered event-causal libertarians a solution to the problem of cross-world luck (see Mele 2006, chap. 5). But I also have argued that agent causationists have basically the same problem. The following passage, which is focused on Timothy O'Connor's agent-causal view, bears on the latter claim. The featured agent, Tim (unlike Jim), has agent-causal power. "Choice"—rather than "decision"—is a key term in the passage. I use the terms interchangeably when what is at issue are decisions and choices to act.

Assume that Tim chose freely in the scenario under consideration. Then, on O'Connor's view, Tim "had the power to choose to continue working or to choose to stop, where this is a power to cause either of these mental occurrences. That capacity was exercised at t in a particular way (in choosing to continue working), allowing us to say truthfully that Tim at time t causally determined his own choice to continue working" ([O'Connor] 2000, p. 74). Suppose that the position reported in the preceding two sentences is true. Why should we suppose that the following cross-world difference is not a matter of chance or luck: that Tim exercised the capacity at issue at t in choosing to continue working rather than in choosing to do something else, as he does in some possible worlds with the same past and laws of nature? Grant that Tim "causally determined his own choice to continue working." Why aren't the differences in his causal determinings at t across worlds with the same past and laws of nature a matter of chance or luck? Tim was able to causally determine each of several choices, whereas a counterpart who fits the event-causal libertarian's picture was able to make—but not to causally determine—each of several choices. If it is a matter of chance that the latter agent chooses to keep working rather than choosing to do something else, why is it not a matter of chance that the former agent causally determines the choice he causally determines rather than causally determining a choice to do something else? (Mele 2006, pp. 54–55)

Notice that the central question here is about the agent's causally determining one choice rather than causally determining another choice. (Readers

should not assume that I regard myself as having a good understanding of what O'Connor means when he says that an agent "causally determined" his choice.) Someone may claim that the cross-world difference between Tim's choosing at t to continue working and his choosing at t to take a break (both worlds having the same past up to t and the same laws) is just a matter of luck, and an agent causationist may reject this claim on the grounds that, in both worlds, Tim causally determines his choice.⁴ But this just moves the debate back a step and opens the door to the central question in the passage just quoted—the question whether the difference at t between Tim's *causally determining* one choice and his *causally determining* the other choice is just a matter of luck.

The next paragraph in Mele 2006 reads as follows:

Perhaps O'Connor is thinking that the conceptual relation between control and chance is such that the fact that Tim exercised direct control over which choice he makes answers each of these questions. Should his readers find this thought persuasive? I do not see why. Even if the fact that Tim exercised direct control in choosing to continue working is incompatible with its being just a matter of luck that he chose to continue working, this does not show that a relevant cross-world difference between his exercising direct control "in [this] particular way" ([O'Connor] 2000, p. 74) and his exercising it in choosing to do something else is not just a matter of luck. The reader should bear two points in mind. First, as I explained, it is not just a matter of luck that Tim chose to keep working even on event-causal libertarian views [i.e., even if Tim lacks agent-causal power]. Second, O'Connor does not place cross-world differences in agents' doings out of bounds in the context of free will: in fact, such differences are *featured* in his objection from chance to event-causal libertarians. A third point also is worth making . . . O'Connor's critique of event-causal libertarianism makes it plain that he does not believe that what Kane [an event-causal libertarian] conceives of as an exercise of direct control . . . solves the problem of cross-world luck at the time of action, and there is a parallel worry about exercises of direct agent-causal control. (Mele 2006, p. 55)

The stage is set for the pair of passages from Pereboom 2014 that I said I would reproduce. The first passage concerns "an event of the following type . . . G: A's causing D at t ," where A is an agent and D is a "decision" (p. 51). Pereboom contends that "the crucial control is not exercised by way of" events prior

to G, “but by the agent-as-substance” (p. 52). The passage continues: “If in addition to the events that precede G we hold fixed in *W* and *W*^{*} the agent-as-substance’s exercise of her agent-causal power, D will occur in *W* and not in *W*^{*} but only because the agent-as-substance causes D in *W* but not in *W*^{*}. Thus it wouldn’t appear to be a matter of luck that D occurs in *W*” (p. 52).⁵

Regarding Tim, I asked why we should “suppose that the following cross-world difference is not a matter of chance or luck: that Tim exercised the capacity at issue at *t* in choosing to continue working rather than in choosing to do something else, as he does in some possible worlds with the same past and laws of nature” (Mele 2006, p. 54). Let *W* be a world in which Tim exercises his agent-causal power at *t* in choosing to keep working and *W*^{*} be a world with the same past and laws in which instead Tim exercises his agent-causal power at *t* in choosing to take a break. Then my question is why we should suppose that this difference between *W* and *W*^{*} at *t* in how Tim exercises his agent-causal power is not a matter of luck. (Again, as I see it, if such a difference is a matter of luck, then the agent’s exercising his agent-causal power as he does is partly a matter of luck.) I do not see how Pereboom has answered this question, unless the answer is a stipulation that the agent-causal power is such that exercises of it are never even partly a matter of luck. If that is the answer, maybe we have learned something about settling. Maybe settling whether one will decide to *A* is supposed to be such that it involves no luck. Maybe settling whether one will decide to *A* requires exercising complete (absolute, total) control over what one does or does not decide. Perhaps this is, in Pereboom’s view, “the control required for basic desert moral responsibility” for a decision that agents who lack settling power do not have (Pereboom 2014, p. 32).

Here is the second passage I said I would reproduce:

Suppose Ralph exercises his agent-causal power in causing his decision to move to New York at *t_n*, and the context is appropriately indeterministic. Was Ralph’s exercise of his agent-causal power merely a matter of luck? The libertarian agent causalist argues as follows. Even though in another world with the same laws of nature and the same past up to *t_n* Ralph does not decide to move to New York, his causing this decision at *t_n* is not a matter of luck. For Ralph causes this decision at *t_n*, while in the alternative world he does not at that time cause this decision. The difference between these two agent-causal worlds is not that the causally relevant events resolve in different ways without the agent settling whether the decision occurs, as would be the case given only the resources of event-causal libertarianism. This is the

concern that the luck objection—and the disappearing agent version in particular—highlights, and it does not also undermine the agent-causal libertarian view. Causally relevant events now resolve the way they do because Ralph exercises his agent-causal power in his causing at *tn* of his decision to move to New York. (Pereboom 2014, p. 52)

My guess about what Pereboom might mean by settling gets some support from this passage. Consider the following pair of sentences from it: “Even though in another world with the same laws of nature and the same past up to *tn* Ralph does not decide to move to New York, his causing this decision at *tn* is not a matter of luck. For Ralph causes this decision at *tn*, while in the alternative world he does not at that time cause this decision.” The pair includes the following claim: Ralph’s causing the decision at issue at *tn* is not a matter of luck because he causes this decision at *tn* and does not do so in the alternative world. Pereboom seems to be saying that no agent-causing of a decision (in an indeterministic world) is even partly a matter of luck. And in light of his contrasting luck and control in the way he does in passages quoted above, he may hold that in agent-causing a decision an agent exercises complete control over whether he decides to *A*. Maybe, for Pereboom, settling whether one will decide to *A* requires exercising such control.

Elsewhere, I have asked what it is about agent causation in virtue of which no agent-causing of a decision is ever even partly a matter of luck (Mele 2006, pp. 68–70). Here I want to focus instead on complete control. I begin with a question about complete control that concerns free actions in general rather than free decisions in particular. Is having complete control over whether one will *A* a plausible general requirement for freely *A*-ing? Imagine a professional basketball player who is a superb free-throw shooter. He sinks 90% of his shots from the foul line. He has a lot of control over whether he sinks his free-throw attempts—certainly more than I do over whether I sink mine. The claim that, even so, he cannot freely sink any given free throw he sinks because he lacks complete control over whether he sinks it sets too high a bar for free free-throw sinking. At least, no view of free action that deserves to be taken seriously sets the bar that high.⁶

Someone who agrees with my claim about the bar may offer the following three-part reply. (1) If this player freely sinks a free throw, he does so partly because he freely *tries* to sink it. (2) He had complete control over whether he would try to sink it. (3) If his sinking a free throw counts as a free action, it does so partly because he had complete control over whether he would try to sink it.

Is this reply acceptable? Imagine that the same player, Sam, has an indeterministic neural randomizer installed in his head that gives him a 99% chance of trying to sink any free throw he intends to sink and a 1% chance of temporarily breaking down instead and not trying to do anything at the time. Regarding any free throw he attempts, he apparently lacked complete control over whether he would try to sink it. The claim that, for that reason, he cannot freely sink any of his free throws sets too high a bar.

Adjust the chances of Sam's trying and his breaking down to 90% and 10%, respectively. Sam just now tried to sink a free throw and his execution was perfect. He scored the point. That he had a 10% chance of temporarily breaking down and not even trying to do what he intended to do makes Sam strange. But an argument would be needed to persuade me that this fact ensures that Sam did not freely sink that free throw.

I asked whether having complete control over whether one will *A* is a plausible general requirement for freely *A*-ing. My answer is *no*, and I have offered some support for that answer. But, of course, there are differences between sinking a free throw and deciding. For example, when Sam sinks a free throw under normal conditions, he has an intention to sink it and he tries to sink it. But, as I see it (see chapter 2), when Sam decides to have dinner at McDonald's tonight, he does not have an intention to decide to have dinner at McDonald's tonight, and he does not try to decide to have dinner at McDonald's tonight. (He might have an intention to make a decision about where to have dinner, and he might try to come to a decision on the matter, but this intention and attempt are not what the preceding sentence is about.) Might it be that freely deciding to eat dinner at McDonald's differs from freely sinking a free throw in such a way that even though the latter does not require complete control, the former does?

The question I just asked raises another one. What does it mean to say that a person has complete control over what he will decide? Having complete control over whether one sinks a free throw would seem to require that there be no chance that the following happens: one has an intention to sink a free throw, tries to sink one, and nevertheless fails to sink the shot. Consider the following proposition: (*Pcc*) having complete control over whether one will decide to *A*, in a case in which one ends up deciding to *A*, requires that there be no chance that one has an intention to decide to *A*, tries to decide to *A*, and nevertheless fails to decide to *A*.

The proposition is problematic. As readers of chapter 2 will have noticed, one problem is the view of deciding that it presupposes. If, although we have intentions to decide what to do and although we sometimes try to decide

what to do, we (typically, at least) do not have intentions to decide to A and do not try to decide to A , then reflection on complete control over whether one will sink a free throw is not likely to tell us a lot about complete control over whether one will decide to A . We need to find a way of understanding complete control in the latter connection that does not presuppose that decisions to A issue from intentions (and attempts) to decide to A .

A second problem with Pcc also merits mention. The complete control we are looking for is supposed to allow for agents' having complete control regarding decisions that they are at no time determined (in the "deterministic causation" sense of "determined") to make. These decisions are supposed to be such that there was a chance that the agent would not make them.

A brief review of the discussion so far may be in order. How did we get to this point in the discussion? Here is a brief answer. I asked what Pereboom means by an agent's *settling* whether a decision he proceeds to make will occur. I looked for guidance in some claims he makes about agents with agent-causal power, and I conjectured, on the basis of those claims, that settling whether one will decide to A is a matter of exercising complete control over whether one will decide to A . That led me to ask what it is for an agent to have complete control over whether he will decide to A . So far, I have not found an answer.

8.2. Control and Settling Whether One Will Decide to A

How do we exercise control over whether we will decide to A ? (I use the future tense here because Pereboom does so in his disappearing agent objection.) In attempting to answer this question, one might reflect on things people do who intend to decide what to do about some issue, focusing on things they do that have an effect on what they end up deciding to do. In some cases, we gather information before deciding. How we gather information can have an effect on what we decide. How we reason about our options can also have an effect on what we decide. Seemingly, in the information gathering and reasoning processes, we are exercising some control over what we will decide (in some typical cases, at least). However, I do not think this is the sort of thing Pereboom has in mind when he speaks of "the control required for basic desert moral responsibility" for a decision (2014, p. 32). If, in information gathering and reasoning processes, we exercise some control over what we will decide (and whether we will decide to A), this may be termed "indirect control." And Pereboom seems to have in mind what is sometimes called *direct*

control (2014, p. 52), a topic I briefly discussed in earlier chapters and will return to in chapter 11.

In Pereboom's disappearing agent objection, we are told that if agent causation is not involved in decision-making, then certain things "do not settle whether the decision *will* occur," "no occurrence of antecedent events settles whether the decision *will* occur," "nothing settles whether the decision *will* occur," and "it can't be that the agent or anything about the agent settles whether the decision *will* occur" (2014, p. 32; italics altered). In each of these quotations, the future tense is used. So, as Pereboom understands agent causation, are agents who decide to *A* supposed to settle what they will decide (or whether they will decide to *A*) *before* they decide to *A*?⁷ If so, how would they do that? Would they agent-cause a decision (or intention) to make another decision a bit later—a decision to *A*? (The former decision is a second-order decision.) Would they agent-cause some process or other that issues a bit later in their deciding to *A*? I doubt this is the kind of thing Pereboom has in mind. What he is claiming, I believe, is that when (that is, at the time at which) people without agent-causal power decide to *A*, they do not settle what they decide (nor whether their decision to *A* occurs), and (if agent causation is possible) when (again, at the time at which) people agent-cause decisions to *A*, they do settle what they decide (and whether their decision to *A* occurs). In the former case, an agent decides at *t* to *A* without settling at *t* what he decides at *t* (and without settling whether his decision to *A* occurs), and in the latter case, in agent-causing at *t* a decision to *A*, an agent settles at *t* what he decides at *t* (and whether his decision to *A* occurs).

In the absence of substantial guidance about what it is for an agent to *settle* at *t* what he decides at *t* (or whether he decides to *A*), how should one proceed? There is an interesting study of the effects of ambient odors on behavior. Robert Baron found that "passersby in a large shopping mall were significantly more likely to help a same-sex accomplice (by retrieving a dropped pen or providing change for a dollar) when these helping opportunities took place in the presence of pleasant ambient odors (for example, baking cookies, roasting coffee) than in the absence of such odors" (1997, p. 498). The presence of pleasant odors significantly increased helping behavior. And if the helping was preceded by decisions to help, the presence of pleasant odors significantly increased the number of decisions to help. Obviously, the people who helped did not take the pleasant odor to be a reason to help. The influencing role at issue was played by non-reasons.

Imagine a pair of agents, Agnes and Eve. Agnes has agent-causal power. And Eve is what Pereboom calls an "event-causal libertarian agent" (2014,

p. 32); she does not have agent-causal power. Both are subjects in a study like Baron's, both smell the pleasant odor, and both decide to help. Passersby were randomly assigned to the good-odor group or the control group; it was just a matter of luck that Agnes and Eve were assigned to the former group. I will suppose that it is also true that in both cases the decision was influenced by the pleasant odor and that neither agent had any idea that this was so.

Might Agnes have exercised complete control over whether she decided to help even though the decision was influenced by the pleasant odor in the way described? Suppose that if Agnes had been in the control group, she probably would have decided not to help (and the same is true of Eve). Does that have any bearing on whether Agnes exercised complete control over whether she decided to help in the actual case? And if Agnes did not exercise complete control over whether she decided to help, might she, even so, have settled what she decided (or whether she decided to help)? Given my weak grip on what it is to have or exercise complete control over a decision and on what it is to settle whether one decides (or will decide) to *A*, I am stumped by these questions.

8.3. Decision-Makers

Pereboom seeks to provide guidance on how “to form a positive conception of agents as substance-causes in a way that does not permit reformulation in terms of agents as causes solely by virtue of their involvement in events” (2014, p. 56). He finds in “the Stoic theory of agency . . . a conception of an agent as having the executive power to determine which of her motivational and doxastic states will result in action” (p. 57). He writes: “It is at least initially intuitive to think that what possesses and exercises this executive power is the agent herself, and not merely the agent's states, or else agent-involving events” (p. 57). Pereboom reports that “in the Stoic theory, in decision and action the agent has an independence of the causal efficacy of all such motivational and doxastic states” (p. 57). “A further feature of the Stoic theory,” he adds, “is that in order for a decision to take place, the agent indeed *must* exercise such executive control. This idea is intuitive. With only the causal efficacy of the various motivational states in place, we don't yet have a decision. Rather, a decision comes about only when the agent makes up his mind and makes it happen” (p. 57).

Although Pereboom says that his disappearing agent objection does not target “agency” itself and rather targets basic desert moral responsibility (2014, p. 32), he here seems to be suggesting that deciding itself depends on

agent causation. Proponents of event-causal views of decision-making will complain about the characterization of their view here. According to such views, it is agents who decide. Agents who decide (at least when they decide) are able to decide and have the power to decide. Proponents of event-causal theories of decision-making do not view an “agent’s states, or else agent-involving events” as exercising the power or ability to make decisions. This is as it should be. States and events are not able to make decisions and therefore are not able to exercise the power or ability to make decisions. Only agents can decide what to do.

A comparison with actions of another kind might help. Recall Sam, the superb free-throw shooter. He has and exercises the ability or power to sink free throws, and he sinks many of them. His intentions, beliefs, skills, and the like do not sink free throws—alone or in combination with one another. And that is no surprise, because they are not able to sink free throws. Does one have to be an agent causationist to say these things about free throws? Not according to proponents of event-causal theories of action. And the same goes for decision-making. It is agents who make decisions, we say. When they do, we say that such things as intentions, beliefs, and desires (or their physical realizers, or facts about what agents intend, believe, and desire) are among the causes of the decision.⁸ But we do not say that things such as these decide or make decisions. We ascribe the decision-making to the agent, just as we ascribe the free-throw sinking to Sam.

Some philosophers have claimed that proponents of event-causal theories of action cannot accommodate actions of any kind. Thomas Nagel writes: “The essential source of the problem is a view of persons and their actions as part of the order of nature, causally determined or not. That conception, if pressed, leads to the feeling that we are not agents at all . . . *My doing* of an act—or the doing of an act by someone else—seems to disappear when we think of the world objectively. There seems no room for agency in [such] a world . . . there is only what happens” (1986, pp. 110–11). David Velleman expresses the worry as follows: “reasons cause an intention, and an intention causes bodily movements, but nobody—that is no person—*does* anything. Psychological and physiological events take place inside a person, but the person serves merely as the arena for these events: he takes no active part” (1992, p. 461). In Mele 2003a, I discussed this worry under the rubric “the problem of disappearing agents” (p. 215).

Suppose, for the sake of argument, that the problem just described cannot be solved by any event-causal theory of action, and suppose, in addition, that action depends on agent causation. What implications do these suppositions

have for the project of Pereboom's 2014 book? As Pereboom notes, he defends "the optimistic view that conceiving of life without" the type of free will that he is skeptical about "would not be devastating to our conceptions of agency, morality, and meaning in life" (2014, p. 4). His positive view is meant to preserve agency, among other things; and to do that, it must preserve actions. However, if action depends on agent causation, an important part of Pereboom's position on agent causation is a potential source of trouble for the optimistic view. A major plank in Pereboom's argument against basic desert moral responsibility and, as he puts it, "the sort of free will required for this sort of moral responsibility" (p. 4) is his argument for the thesis that "it's doubtful for empirical reasons . . . that we are agents of the sort specified by" agent-causal libertarianism (pp. 69–70). If action itself depends on our being agents of that sort, then if we are not agents of that sort, we never act.

There are options here. For example, Pereboom can claim that although deciding requires agent causation, lots of important human actions do not. Another option is to claim that although the exercises of agent-causal power to which agent-causal libertarian theories appeal probably do not exist, other exercises of agent-causal power do exist and support the occurrence of actions (though not free actions). These options are interesting, and Pereboom mentions that "deterministic agent-causal theory of action is available to the free will skeptic" (2014, p. 105).⁹ But he does not argue for any of these options in his 2014 book. When he introduces his disappearing agent objection, he says that it "targets basic desert moral responsibility rather than agency" (2014, p. 32). This claim is associated with another option—leaving it open that deciding does not depend on agent causation. That clearly is left open in Pereboom's formulation of his disappearing agent objection, as quoted in my introduction to this chapter, and I treat it as open in the remainder of this chapter.

I close this section with a response to a suggestion made by a referee, Derk Pereboom. He suggested that what he means by "settling" is what I mean by it when I assert that in deciding to *A* one settles on *A*-ing. In my view (as in Pereboom's), decision-making is compatible with determinism, but attention can be focused now on decision-making in indeterministic scenarios in which what an agent decides at *t* is open all the way up to *t*. In this kind of scenario, Pereboom's suggestion (qua referee), more fully stated, is that when one is faced with, for example, a pair of options, *A* and *B*, one settles whether one will be in a decision state of being settled on *A*-ing or a decision state of being settled on *B*-ing simply by deciding to *A* or deciding to *B*. Now, in my view, an agent's deciding to *A*—in general and in a scenario like this—does

not depend on his having agent-causal powers and is accommodated by an event-causal libertarian view. Call this thesis *T*. Suppose that *T* is true. Then if an agent's deciding to *A* in a case of the kind in question is sufficient for his settling whether he is in one decision state or another, such settling does not depend on agent-causal powers and the disappearing agent objection fails. So suppose that *T* is false and, more specifically, that both conjuncts are false. Then if Pereboom's skepticism about agent causation is justified (that is, the specific skepticism about it commented on in this chapter), we should believe not only that we never decide freely but also that we never decide at all in indeterministic cases of the sort at issue. Because this result—that we never decide at all in these scenarios—is at odds with Pereboom's report that his disappearing agent objection does not target agency (2014, p. 32), I infer that the notion of settling at work in my view of decision-making is not the notion of settling at work in his disappearing agent objection.

Another observation on this matter is in order. If, in situations of the sort at issue, deciding to *A* is sufficient for settling whether one is in a decision state of being settled on *A*-ing or a decision state of being settled on *B*-ing, perhaps Pereboom's disappearing agent objection can be reformulated in a way that features deciding and does not mention settling at all. What follows is such a reformulation that preserves most of the original wording:

The modified disappearing agent objection: Consider an intention that arises in a context in which the agent's moral motivations favor that intention, and his prudential motivations favor his coming to have a contrary intention instead, and the strengths of these motivations are in equipoise. On an event-causal libertarian picture, the relevant causal conditions antecedent to that intention's arising—that is, the occurrence of certain agent-involving events—do not enable the agent to decide what to do, but only render the occurrence of the intention about 50% probable. In fact, because no occurrence of antecedent events enables the agent to decide what to do, and only antecedent events are causally relevant, *nothing* enables the agent to decide what to do. Thus the agent can't decide what to do, and he therefore lacks the control required for basic desert moral responsibility for his intention.

An argument for the crucial proposition that the agent is not able to decide what to do, given relevant antecedent events, is conspicuously absent here. Why is an "event-causal libertarian agent" (Pereboom 2014, p. 32) who is in the situation described here unable to decide what to do? What is the

argument that takes us from the description of the case to the conclusion that the agent is unable to decide what to do? If it is claimed that the agent is unable to decide what to do because he cannot settle what decision state he enters and deciding what to do requires such settling, we need an interpretation of settling that makes it clear why this claim is supposed to be true. My account of deciding (see chapter 2) does not yield such an interpretation. On my account, there is no more to settling on *A*-ing—and, in so doing, entering a state of being decided upon *A*-ing, a decision state—than there is to deciding to *A*. And, on my account, there is no more to settling what decision state one enters, when one is faced with live options, than there is to deciding to *A* when one is faced with such options.

8.4. Settling and Complete Control

Readers who are attracted to libertarianism may have a variety of different attitudes toward agent causation. Here are three among many. (1) They may deem it impossible. (2) They may be agnostic about whether it is possible and believe, with Pereboom (2014, chap. 3), that the balance of evidence supports the claim that agent causation is not actual. (3) They may not see how it is an improvement on event-causal libertarianism (Balaguer 2014, pp. 84–85; Haji 2004; Mele 2006). If such readers assume that settling whether one decides to *A* (whatever, exactly, such settling may be) requires agent-causal power, they may reject the idea that freely deciding to *A* and basic desert moral responsibility for a decision to *A* require settling whether one decides to *A* (see Palmer 2013, p. 110). If the same readers assume that settling whether one decides to *A* depends on agent-causal power because such settling requires exercising complete control over what one decides, they may reject the idea that freely deciding to *A* and basic desert moral responsibility for a decision to *A* require exercising complete control over what one decides.

A key claim in Pereboom's disappearing agent objection is that if neither the agent nor "anything about the agent settles whether the decision will occur," the agent lacks "the control required for basic desert moral responsibility for it" (p. 32). I do not understand this claim well enough to assess it, because, despite my efforts, I do not have a good grip on what is meant by settling whether a decision will occur.

An article by Meghan Griffith (2010) may be thought to fill the gap left by Pereboom. Central to her reply to the problem of present luck is the following claim, which Griffith dubs RUL: "In the free will context, something is a matter of luck for *A* if and only if it happens to *A*" (p. 46). She writes: "RUL

cashes in on the distinction between doing and happening. What happens to someone is not something she does. In the free will context, that something happens to *A* is sufficient for its being a matter of luck, since it rules out her having done it. Since this is ruled out, her responsibility for it is undermined” (p. 46). In Griffith’s view, decisions involving exercises of agent-causal power do not happen to the agents and therefore are not even partly a matter of luck. Her view includes a position on agent-causal control that I will discuss shortly.

As readers of chapter 5 will have surmised, I believe that RUL is false. Distinguish trying to vote for Gore from voting for Gore. Al does both in my story about him in chapter 4. His doing the latter was partly a matter of luck. But it was not something that happened to him. Al’s voting for Gore was an action—and a free action, I suggested, if free actions are common in his world. It was—luckily—a successful attempt to vote for Gore.

Here is another example. Ann is a hockey player. She takes a shot at the goal. The puck is veering a bit wide, when it ricochets off a defender and bounces into the goal. Ann scored a goal. That’s something she did, not merely something that happened to her. And her scoring the goal involved some luck; it was partly a matter of luck. Furthermore, if Ann is a normal, rational unmanipulated agent in normal circumstances and free actions are common in her world, it is plausible that her scoring the goal was a free action.

If Al’s voting for Gore and Ann’s scoring the goal are free actions, they may be deemed indirectly free (see chapter 5, section 5.1). Griffith may seek to avoid the problem I identified for RUL by modifying it as follows: (RUL*) In the context of directly free actions, something is a matter of luck for *A* if and only if it happens to *A*. Now, I claimed that Bob’s deciding to cheat in my story about him in chapter 6 is partly a matter of luck. In my view, again, to decide to do something is to perform an action of a certain kind. And, on Griffith’s view, the very fact that Bob performed the action of deciding to cheat is incompatible with that action’s being partly a matter of luck. In her view, apparently, *x* is not even partly a matter of luck for a person unless *x* happens to that person; and, as I have mentioned, she asserts that the fact that *x* happens to a person “rules out her having done it” (2010, p. 46).

Griffith’s view does not fare well when it comes to such actions as Al’s voting for Gore and Ann’s scoring the goal. They are luck-involving actions—actions that are partly a matter of luck. Does it fare any better in the case of Bob’s deciding to cheat?

Griffith says that when we take an event-causal libertarian perspective on an agent faced with a pair of options, *A* and *B*, we see that “it just happens

to her that the decision is to *A* rather than to *B* because she is missing the power to determine the decision” (2010, p. 51). When Bob is viewed from this perspective, Griffith presumably would say that it just happens to him that his decision is to cheat rather than to flip the coin. And, guided by RUL (or RUL*), she would deduce that the fact that his decision is to cheat rather than to flip the coin is a matter of luck. But if that fact is a matter of luck, what recommends the view that his deciding to cheat is not even partly a matter of luck? Griffith can claim that it cannot be even partly a matter of luck because it is an action and therefore is not something that happens to Bob. But this line of defense is unsatisfactory, as I have explained. Ann’s scoring the goal is an action *and* partly a matter of luck, and the same is true of Al’s voting for Gore.

In Griffith’s view, the event-causal libertarian perspective leaves out something that is required for free will—namely, agent causation. Agents, as characterized by event-causal libertarians, are said to be “missing the power to determine the decision” (Griffith 2010, p. 51). What power is that? What does it amount to? What is it for an agent to *determine* a decision? Griffith asserts that an agent who lacks the agent-causal power “seems not to have control over the crucial element for which she is responsible: that she has decided to *A* rather than to *B*” (p. 50). She compares this agent to a man in a story by Robert Kane who tries to smash a glass table by hitting it with his arm (Kane 1999b, p. 227). The table broke, but it was undetermined whether the man’s striking it would break it. Griffith writes: “Although [he] causes the table’s breaking, he does not *completely control* whether the table breaks. In this sense, its breaking happens to him” (p. 50; emphasis altered). Taking our lead from this, we have an answer to my questions about the alleged power to “determine the decision.” It is the power to *completely control* which decision one makes.

What is it about agent causation that underwrites the claim that in scenarios like Bob’s (some) agent-causes have the power to “completely control” which decision they make? A good answer to that question might enable me to see that and why the problem of present luck is an illusion—or a problem only for libertarian views that make no use of agent causation. I have not yet seen such an answer.

In my view, even if the difference between what an agent does at *t* in one world and what he does at *t* in another world with the same past up to *t* and the same laws of nature is just a matter of luck, the agent may perform a directly free action at *t* in both worlds (Mele 2006, chap. 5). This is the thesis I labeled *LDF* in chapter 6. Perhaps, partly because I accept *LDF*, I am not in

search of a notion of control that allows for “complete control” over what one decides to be exercised at t in one or both worlds.

Suppose that the pertinent difference at t between a world in which an agent decides at t to A and a world with the same past up to t and the same laws of nature in which he decides at t to B is just a matter of luck. If I had a good grip on the idea of complete control over whether one decides to A or decides to B , I might find myself believing that the agent does not exercise complete control over whether he decides at t to A or instead decides at t to B . Even so, if *LDF* is true, he may make these decisions freely. (As I have mentioned, by “complete control” I do not mean “as much control as metaphysically possible.” I leave it open that the following conjunction is true: exercising the power of agent causation is required for exercising complete control over whether one decides at t to A or instead decides at t to B in scenarios of the sort at issue, and agent causation—conceived of in such a way as to allow for an agent’s having complete control over what he decides in scenarios of the sort at issue—is metaphysically impossible.)

Timothy O’Connor (2011) quotes the following from Mele 2006, p. 70:

[*L.*] If the question why an agent exercised his agent-causal power at t in deciding to A rather than exercising it at t in any of the alternative ways he does in other possible worlds with the same past and laws of nature is, in principle, unanswerable . . . because there is no fact or truth to be reported in a correct answer . . . and his exercising it at t in so deciding has an effect on how his life goes, I count that as luck for the agent. (2011, pp. 324–25)

He then writes:

Suppose we take this as a stipulative account (or sufficient condition) on luck “as Mele understands the notion.” If so, it is open to the agent causationist to deny that luck in this stipulated sense is of any significance whatsoever—not, for example, being relevant to freedom and moral responsibility. Mele in fact agrees! He does not press his luck objection as a deep skeptical worry about indeterministic freedom . . . Instead, he wields it to neutralize the agent causationist’s objection to causal indeterminism. (2011, p. 325)

My response is yes and no. Yes, I have never appealed to luck in an argument for the falsity of libertarianism. In fact, I have never argued for the falsity of

be examined for plausibility. And someone who believes that there is no such fact or truth has the option of trying to explain why the truth of that belief is compatible with the agent's acting freely at the time. Either way of proceeding might result in progress. I take the latter route, as I have mentioned.

Here, to keep attention focused on the problem of present luck, I repeat a question I raised in chapter 6 about the story of Bob and the coin. Given the details of Bob's story (in either the original version or the "fuller" version), how can Bob have enough control over whether he decides to cheat or does something else instead at noon for his decision to be directly free and for him to be directly morally responsible for it? This, as I observed, is an instance of the central question posed by the problem of present luck. And notice that it neither calls for something to be contrastively explained nor mentions luck.

O'Connor reports that in reply to my worry about present luck, "[Randolph] Clarke (2005) argues that an agent causal capacity would provide a stronger variety of control than is available on causal indeterminism" (O'Connor 2011, p. 325). Be that as it may, recall Clarke's report that, in his judgment, relevant arguments collectively "incline the balance against the possibility of substance causation in general and agent causation in particular" (2003, p. 209). Clarke argues (2003) that agent-causal powers are required for free will (at least, if incompatibilism is true). If agent causation is required for free will and impossible, free will is impossible. Recall the reference to a double-edged sword in chapter 6.

As I have observed, some agent causationists regard the problem of present luck or something very similar as a decisive problem for event-causal libertarianism. And yet, some of the same agent causationists seem to regard the problem as no threat at all to agent-causal libertarianism. Why might that be? Consider the following from O'Connor: "The agent causationist takes it to be a virtue of her theory that it enables her to avoid a 'problem of luck' facing other indeterministic accounts. Agent causation is precisely the power to directly determine which of several possibilities is realized on a given occasion" (2011, p. 325). This claim may be combined with the idea, mentioned earlier, that this determining power is the power to completely control which decision one makes to yield the following assertion: Agent causation is precisely the power to completely control which decision one makes.

Now, anyone can *say* that something or other is the power to completely control which decision one makes. An event-causal libertarian can say this about some non-agent-causal decision-making power, and so can a noncausalist libertarian. One thing I would like to know is how replacing event-caused decisions or uncaused decisions with agent-caused decisions (or intentions) is

supposed to make true something that would otherwise supposedly be false—namely, that the decisions are made by someone who exercised the power to completely control which decision he made (or which intention he came to have). (On this sort of thing, see Mele 2006, chap. 3.) A plausible answer would be a plausible reply to the problem of present luck. Until I see such an answer, I will take comfort in *LDF*.

O'Connor (2011, p. 311) notes that Richard Taylor (1966) “propounded agent causation as a feature of all intentional action.” It is implausible (to put it mildly) that in the case of every intentional action, agents have and exercise complete control over whether they perform that action. When a professional basketball player with a 90% success rate at sinking free throws sinks a free throw in the normal (for him) way, he intentionally sinks it. But he would seem not to have complete control over whether he sinks it. A free-throw shooter who has and always exercises complete control over whether he sinks his free throws sinks every free throw he tries to sink. So, as at least some agent causationists think of agent causation, the agent-causal power to *A* seemingly does not include complete control over whether one *A*-s as an essential feature (for all *A*-s). And if having “the power to directly determine which of several possibilities is realized on a given occasion” (O'Connor 2011, p. 325) necessarily includes having complete control over which of the possibilities is realized on that occasion, then, as some agent causationists conceive of agent causation, the agent-causal power to *A* seemingly does not always include the power to directly determine that one *A*-s.

A few paragraphs ago, I asked what it is about agent causation that underwrites the claim that in scenarios like Bob's (some) agent-causes have the power to “completely control” which decision they make. One option for someone who conceives of agent causation differently than Taylor does is simply to stipulate that agent causation is the power to completely control which decision one makes or which intention one comes to have in situations of the sort at issue. But the stipulation is unilluminating.

What is complete control? And what is it for an agent to determine something? Considering these questions in the context of a simple real-world experiment will prove useful. Only part of the experiment is relevant for my purposes. That is the part I will describe.

Subjects are instructed to press either the Q key on a computer keyboard with their left index finger or the P key with their right index finger (and never to press both at the same time). They are told that which key they press is up to them. They will make over forty key presses—either Q or P, sometimes one and sometimes the other—in the course of an hour, and they are

asked to refrain from planning in advance which key to press. Before they press they will hear a tone that they are instructed to treat as a “decide” signal. When they hear the signal, they are to decide right then which key to press and then press it straightaway. Pressing a key, as defined for the purposes of this experiment, requires that the key move all the way down and make contact with the switch under it.

What would it be for Sol to have complete control over whether he presses Q or P when he hears the next “decide” signal? Consider the following suggestion. For Sol to have complete control over this is for the following things to be true: at the relevant time Sol is able to try to press Q and able to try to press P, and, regarding each key, if he tries to press it there is no chance that he will fail to press it.

This suggestion leaves out something important. Doesn’t Sol’s having complete control over which key he presses require his having complete control over which key he *tries* to press? What might that amount to? Here is a suggestion to consider. For Sol to have complete control at the time over which key he tries to press is for the following things to be true: regarding each key, when Sol hears the tone, he is able to decide to press it right then and, regarding each key, if he decides to press it right then, there is no chance that he will fail to try to press it.

There is a predictable worry about this suggestion too, of course. Doesn’t Sol’s having complete control in this scenario over which key he tries to press depend on his having complete control over which key he *decides* to press. What does that amount to?

When faced with my question about complete control over key presses, I looked to trying for an answer. And when faced with a parallel question about trying, I looked to deciding for an answer. Now that the question is what it is for Sol to have complete control over which key he decides to press, where should I turn?

Both of the suggestions I considered have a “no chance” clause. The reason for this is obvious. Suppose that the keyboard Sol is using has a randomizer on it that ensures that there is always a small chance that a key he is trying to press will stick and fail to make contact with the switch under it. (Recall the definition of a key press above.) Then Sol never has complete control over whether he presses the Q key or the P key. Suppose now that a randomizer has been installed in Sol’s brain that ensures that there is always a small chance that his proximal decisions to press a key—his decisions to press a specific key straightaway—will not be followed by a corresponding attempt. Then Sol never has complete control over whether he tries to press the Q key or tries to

press the P key (at least when no alternative route to attempt-production is in use—that is, no route that does not include decisions).

These observations prompt the following two questions: Might Sol nevertheless freely have pressed the Q key the last time he pressed it? Does a satisfactory account of a person's having complete control over whether he decides to *A* or decides to *B* have a "no chance" clause?

The correct answer to the first question, I believe, is yes, provided that free actions are common in Sol's world and it is possible for agents to act freely when their options are of the kind featured in Buridan's ass scenarios. To see why, compare Sol's pressing the Q key with Al's voting for Gore or Ann's scoring the goal.

What about the second question? According to a standard libertarian view, if a person's deciding to *A* is to be a directly free action, there was a chance, right up to the time at which the decision was made, that he would not decide then to *A*. But this alone does not obviously commit a proponent of this view to claiming that no satisfactory account of a person's having complete control over whether he decides to *A* or decides to *B* can have a "no chance" clause. One reason is that the following option is not obviously a nonstarter: having complete control over whether one decides to *A* or decides to *B* is *not required* for a directly free decision in favor of one or the other of these courses of action.

Assessment of this option would benefit from an acceptable account of what it is to have complete control over whether one decides to *A* or decides to *B*. It can be said that having the kind of control at issue is a matter of its being *entirely up to* the agent whether he decides to *A* or decides to *B* or a matter of the agent's having the power to *determine* whether he decides (or intends) to *A* or decides (or intends) to *B*. But it is not as though the key terms here carry their meanings on their faces. Just as I would like to be told what it is to have complete control over whether one decides to *A* or decides to *B*, I would like to be told what it is for it to be entirely up to an agent whether he decides to *A* or decides to *B* and what it is to have the determining power at issue. From my point of view, it would be wonderful if the initial answers to my questions did not mention agent causation and then someone explained why agent causation is supposed to be needed for complete satisfaction of the conditions identified in those answers. But that wonderful state of affairs might not be in the cards. Owing to my ignorance about how I am supposed to understand the key expressions (for example, "complete control over"—or its being "entirely up to one"—whether one decides to *A* or decides to *B*), I do not know whether the (alleged) phenomena they are supposed to pick out

can be given informative preliminary characterizations that do not appeal to agent causation.

Return to Sol. Imagine now that one thing that can happen when he detects the “decide” signal is that his ability to make decisions is temporarily eliminated. Sol is paid one dollar for each button press; so the temporary loss of the ability has a down side. When the signal is emitted, the following things are true by hypothesis: Sol is able to respond to it with a decision to press Q; Sol is able to respond to it with a decision to press P; there is a chance that Sol will very soon temporarily lose his ability to make decisions. (Bear in mind that it takes several milliseconds to detect the signal.) Consider a trio of possible worlds in which everything is the same up to the time the “decide” signal is emitted. In W_1 , Sol proceeds to decide to press Q; he makes this decision at t . In W_2 , he decides at t to press P. In W_3 , at t , his decision-making ability is temporarily eliminated. When he detected the signal, did Sol have complete control over whether he would decide to press Q or decide to press P?

Compare W_1 with W_3 . The difference between them at t would seem to be just a matter of luck, and the same goes for the difference at t between W_2 and W_3 . But if these differences are just a matter of luck, are not Sol’s deciding to press Q and his deciding to press P partly a matter of luck? If Sol had had the bad luck in those worlds that he had in W_3 , he would not have made either decision. Someone may assert that even if these decisions are partly a matter of luck, Sol had complete control in W_1 and W_2 over whether he decided to press P or decided to press Q. The claim has the ring of a falsehood to me. But because I do not know how someone who might make this claim might intend “complete control” over this to be understood, I would not reject the assertion before asking for clarification. If the truth of claim (C_1) that Sol did not in W_1 have complete control over whether he decided to press P or decided to press Q cannot be derived from the truth of the claim (C_2) that, in W_1 , Sol’s decision at t to press Q was partly a matter of luck, then perhaps the following claim cannot be derived from C_2 either: (C_3) Sol did not (directly) freely make the decision at issue. And if that is so, then given that I am still in the dark about what the complete control mentioned in C_1 is supposed to be, it makes sense for me to continue to take comfort in *LDF* and continue to offer my own solution to the problem of present luck (Mele 2006, chap. 5).

My solution—which makes no appeal to agent causation—has the virtue of being fairly easy to understand. In that respect, at least, it has an advantage over proposed agent-causal solutions to the problem of present luck that leave us in the dark about how such key terms as “complete control” and “determine the decision” are to be understood.

8.5. Deciding Freely?

The present section provides some motivation for the thought that an agent like Eve—an “event-causal libertarian agent” (Pereboom 2014, p. 32)—can decide freely and have basic desert moral responsibility for her decision. I am not hoping to move Pereboom or Griffith toward embracing this thought. My aim is to sketch a picture for readers in general.¹¹

Suppose that Eve is sane, rational, highly intelligent, uncompelled, uncoerced, unmanipulated, and very well-informed about the issue that is the topic of her current deliberation. The practical question she faces is whether to accept or reject a job offer, and she has carefully gathered relevant information and painstakingly assessed pros and cons. At no point is it determined what Eve will decide. She is able to decide to accept the offer for reasons that recommend accepting it and able to decide to reject it for reasons that recommend rejecting it. There are many good reasons on both sides, and either decision would be reasonable. In the end, Eve decides to reject the offer; she rationally makes that decision—at t —for reasons that recommend rejecting it.

In another possible world with the same past up to t and the same laws of nature, Eve rationally decides at t to accept the job. She makes that decision for reasons that recommend accepting it. Someone who is not finicky about usage of the word “luck” can reasonably say that the difference I mentioned between the two worlds at t is just a matter of luck and that Eve’s decision, therefore, is partly a matter of luck. Depending on how settling what one decides (or will decide) is supposed to be understood, Eve’s decision’s being partly a matter of luck may preclude her having settled what she decided (or what she would decide). But does it preclude her having decided freely and her having basic desert moral responsibility for her decision?

There was some chance that Eve would decide at t to reject the job and some chance that she would decide then to accept it. It is likely that there was also a chance that Eve would, at t , continue to be unsettled about what to do and continue deliberating about the job offer, and there might have been a chance of her mind beginning to wander at t (among other things). Consider these chances—or antecedent probabilities—very shortly before t . If they came out of the blue, wholly as a matter of luck, that would be cause for worry about Eve. But they did not. They were shaped in significant part by Eve’s evidence gathering and her deliberation. These antecedent probabilities were also shaped partly by Eve’s long-term preferences and values, and these preferences and values were in turn shaped by past decisions Eve made, by what she learned from past successes and mistakes in decision-making, and so on.

Imagine that we are like Eve in that, sometimes, the processes that issue in our decisions are such that it is at no time determined what we will decide. Even so, our values, preferences, learning history, information gathering, deliberation, and so on constrain the physically and psychologically possible outcomes and shape the antecedent probabilities of the outcomes. If we come to know that we are like Eve in this respect and the indeterminism worries us, we should do our best to minimize our chances of making poor decisions by working on developing good habits of decision-making and good habits in general. If we learn that we are like Eve in the respect at issue, should we also infer that we never act freely?

It may be said that someone like me who believes that Eve may decide freely and have basic desert moral responsibility for her decision sets the bar for free action and this sort of moral responsibility very close to where a compatibilist sets it. This might make me worry, if there were an argument that convinced me that compatibilism is false. But there is no such argument. I have argued elsewhere (against arguments to the contrary) that both compatibilist and event-causal libertarian views are live options (Mele 1995a, 2006), and I have not yet seen convincing grounds for rejecting that thesis.

8.6. Parting Remarks

Pereboom places his disappearing agent objection in the family of “luck objections” (2014, p. 32). He reports that, in his view, it is the member of this family of objections “that reveals the deepest problem for event-causal libertarianism” (p. 32). Now, I myself have articulated a luck-featuring problem for standard libertarian views (Mele 2006, chap. 3); and, as I mentioned, I offered event-causal libertarians a solution to the problem in the same book (chap. 5). Owing to my inadequate grip on what Pereboom means by an agent’s settling whether a decision will occur (and on what it is to exercise complete control over what one will decide, if that is what he means), I am not in a position to assess the proposal that the problem posed by his disappearing agent objection is deeper than, for example, what I called “the problem of present luck” (Mele 2006, p. 69).

An ideal, reader-friendly version of Pereboom’s presentation of his disappearing agent objection would include substantial guidance on what it is for an agent to settle whether a decision will (or does) occur. Ideally, the reader would get an initial sketch of an account of this settling that does not refer to agent causation and then an explanation of why it is that “event-causal

libertarian agents” cannot do anything that satisfies the account. If no sketch of decision settling is even intelligible without reference to agent causation, a reader-friendly version of the presentation would explain why that is so. Also if the guidance provided on settling does not render it obvious that basic desert moral responsibility for a decision one made depends on one’s having settled whether that decision would occur, a reader-friendly version of the presentation would include an argument for this thesis about settling and responsibility. Such an argument may prove to be very interesting.

I conjectured that, for Pereboom, settling whether one decides (or will decide) to *A* is a matter of exercising complete control over whether one decides (or will decide) to *A*. If my conjecture is on target (or even if it is not), it would be great to have an analysis of “*S* exercises complete control over whether he decides to *A*.” That is a lot to ask for, I realize; but substantial steps in that direction would be appropriate and useful.

As I have explained, I am unpersuaded by Pereboom’s disappearing agent objection to event-causal libertarianism. I continue to see event-causal libertarianism as a legitimate contender (along with compatibilism). However, it would be good to see a version of Pereboom’s disappearing agent objection that has the reader-friendly features I just mentioned. I do not take it for granted that no version of this objection will succeed. An account of what it is to settle whether one will (or does) decide to *A* may prove illuminating. Possibly, with an account of such settling in place, we will be able to ascertain, among other things, whether settling whether one will (or does) decide to *A* is required for freely deciding to *A* and for basic desert moral responsibility for one’s decision.¹²

Notes

1. I am not alone in this. See Palmer 2013, pp. 107–12 on earlier versions of Pereboom’s disappearing agent objection.
2. As I understand (2), the idea is not that one decides to *A* and then immediately enters a state of being decided upon *A*-ing. Instead, one enters that state at the very time at which one decides to *A*.
3. The general problem applies to other actions too, but I focus on decisions here.
4. Recall my convention about counterparts. See chapter 4, n. 11.
5. I mentioned that “decision” sometimes refers to the act of deciding and sometimes to “the immediate issue of the act, a decision *state*, a state of being decided upon something” (Mele 1992a, p. 158). Which is D supposed to be? Pereboom writes: “What the agent-causal libertarian posits is an agent who possesses a causal

power, fundamentally as a substance, to cause a decision—or more comprehensively, as O'Connor (2009) specifies, 'the coming to be of a state of intention to carry out some act'—without being causally determined to do so, and thereby to settle, with the requisite control, whether this state of intention will occur" (2014, p. 51). A "state of intention" to *A* formed in an act of deciding to *A* is a decision state.

6. Readers who believe that free actions must have moral significance should feel free to augment my free-throw stories with the detail that the player was offered a bribe to miss a relevant shot. Similar additions may be made to other sporting examples of mine.
7. I include the parenthetical clause because "settle whether" is an operative expression in the quotations from Pereboom here.
8. For discussion of the disjunction here, see Mele 2013a.
9. This idea is developed in Pereboom 2015. I mentioned Pereboom's uncertainty about the conceptual and metaphysical possibility of agent causation in his 2014 book. He writes: "It may turn out that fundamental substance causation is metaphysically impossible, or even conceptually impossible" (2014, p. 58). This claim is not restricted to substance causation (including agent causation) in an indeterministic setting.
10. I am grateful to Michael Robinson for encouraging me to comment on this point.
11. The ideas I am about to sketch are much more fully developed in Mele 2006, chap. 5. See also Mele 2013d, pp. 248–55.
12. I am grateful to Randy Clarke and Derk Pereboom for comments on a draft of Mele n.d.b, on which much of this chapter is based, and to Stephen Kearns and an audience at the University of Calgary (March, 2015) for discussion. Section 8.4 derives from Mele 2014b, and I am grateful to Meghan Griffith for her comments on a draft of that article.

Libertarianism and Human Agency

SOME SCIENTISTS HAVE reported what they regard as evidence of indeterministic brain processes that influence behavior (Brembs 2011; Maye et al. 2007). How do these reports bear on the positive side of libertarianism about free will? That is an approximation of my guiding question in this chapter. I make the question more precise in section 9.1, in light of some conceptual and scientific background. In the remainder of the chapter, I seek—and eventually offer—an answer.

9.1. My Question and Some Background

Recall that libertarianism about free will is the conjunction of two theses:

1. *The incompatibility thesis*: Free will is incompatible with determinism. In terms of possible worlds, in any possible world in which determinism is true, there is no free will.
2. *The pro-free-will thesis*: There are actions that are or involve exercises of free will—*free actions*, for short.

As I observed in chapter 7, typical libertarians also endorse a more specific version of 2, namely, *LFT*: There are indeterministic agents who sometimes act freely when their actions are not deterministically caused by proximal causes. In chapter 1, I pointed out that different libertarians take different positions on what free will requires beyond the falsity of determinism. Some contend that only beings with agent-causal powers can have free will. Others argue that only uncaused actions can be free. Event-causal libertarians avoid appealing to agent causation, and they typically claim that paradigmatically free actions are indeterministically caused by their proximal causes (Kane 1996).

My guiding question is roughly this: If incompatibilism is true and if the empirical claim and hypotheses I am about to discuss are true, what contribution might indeterministic agent-internal processes of the kind at issue make to free will beyond being sufficient for the falsity of determinism? The empirical claim and hypotheses to be discussed provide no support specifically for the existence of agent-causal powers (as opposed to support for something that a being's having such powers is often regarded as requiring—namely, the falsity of determinism) and no support for the occurrence of uncaused actions.¹ So I focus on event-causal libertarianism.

I turn now to a scientific claim. If one is seeking hard evidence of indeterministic brain processes in animals, it makes sense to start small. Small, relatively simple brains are easier to study than large, complicated ones. Alexander Maye and colleagues report what they regard as evidence of indeterministic brain processes in fruit flies that affect the flies' behavior (Maye et al. 2007; also see Brembs 2011). They also offer two (mutually compatible) hypotheses about why indeterministic brain processes for behavior initiation might have evolved. One is that unpredictability is required for survival (Maye et al., p. 8), and indeterministic brain processes of the sort at issue would result in unpredictable behavior. A predictable pattern of response to pursuit or attack tends to make an animal relatively easy prey (Brembs 2011). The other is that an animal's unpredictable behavior enables it to learn "which portions of the incoming sensory stream are under operant control by [its] behavior" (Maye et al., p. 8). To be sure, it is possible for an animal to be unpredictable to potential predators and to itself even in a deterministic universe. Even so, if animal brains are indeterministic organs, as Maye et al. argue, the evolution of adaptive unpredictability might have benefited from that feature of brains.

If even one actual brain is an indeterministic organ, then determinism is false and a necessary condition for the existence of free will, as libertarians conceive of it, is satisfied. But, of course, it is a long way from indeterministic behavior-production to free will. Presumably, fruit flies lack free will even if some of their behavior is produced by indeterministic brain processes. This observation takes us back to my guiding question.

9.2. Event-Causal Libertarianism and Compatibilism

Some philosophers have distinguished between what may be termed *late* and *early* indeterministic processes in an action-producing stream (Dennett 1978, pp. 294–95; Mele 1995a, p. 212, and 2006, pp. 112–14). Late processes of this

kind are still at work when the actions they issue in begin, and early processes are not. An indeterministic process that generates mental representations of options and stops some time before any option is selected is an example of an early process. An example of a late process is an indeterministic decision-producing process that does not end before it indeterministically issues in a decision to *A*.

I invite readers to imagine that the following two propositions are true. First, indeterministic agent-internal processes that play a role in producing behavior are part of our evolutionary heritage and some such processes evolved because of their contribution to survival-promoting unpredictability. Second, some of the indeterministic processes at work in us are late processes and they include processes that indeterministically issue in decisions. Some readers may be curious about the low-level mechanics of indeterministic agent-internal processes. I do not explore that issue here, but I mention one alleged possibility—namely, that there are quantum probability clouds associated with calcium ions moving toward nerve terminals (Stapp 2007, pp. 30–32).

What might a kind of control—a kind of *regulative* control, if you like this label—that depends on late indeterministic agent-internal processes do for libertarians that deterministic control (or guidance control) cannot do for them? (On regulative and guidance control, see chapter 7, section 3.) Obviously, the existence of the former sort of control—unlike the latter—is incompatible with the truth of determinism and therefore sufficient for the satisfaction of a necessary condition for free action, if incompatibilism is true. But there is more. Many libertarians hold that, necessarily, a being performs a directly free action, *A*, only if, at the time of action, he could have done otherwise than *A* in a sense of “could have done otherwise” that requires the falsity of determinism.² In terms of possible worlds, the claim is this: (*FAP*) Necessarily, a being who *A*-s at *t* *A*-s directly freely only if there is another possible world with the same past up to *t* and the same laws of nature in which, at *t*, he does not *A* and does something else instead. Some late indeterministic agent-internal processes might contribute to the existence of alternative possibilities of this kind; and, of course, late deterministic agent-internal processes cannot do this.

An incompatibilist may be persuaded by a Frankfurt-style case that what Harry Frankfurt (1969) called the “principle of alternate possibilities” (*PAP*) is false (see Pereboom 2001; Stump 1990; Stump and Kretzmann 1991; Zagzebski 1991).³ Such an incompatibilist might also be persuaded by such a case that *FAP* is false. As he or she sees things, the alleged benefit just

mentioned of indeterministic control might not amount to much. This issue merits attention.

In a Frankfurt-style case, if it hits its target, an agent is morally responsible for, say, deciding to steal a certain car even though he could not have done otherwise than that at the time. The agent decides on his own at t to steal the car; but if that had not happened, he would have been compelled to decide at t to steal it. As I have explained elsewhere (Mele 2006, chap. 4), it is open to a libertarian to accept that some indeterministic Frankfurt-style cases do hit their mark and falsify various alternative-possibility principles both about moral responsibility and about free action. Such a libertarian may reject *FAP* and accept a variant of it that requires for directly free A -ing, not that the agent could have done otherwise than A at the time, but instead that the proximal causes of his A -ing indeterministically cause it (Mele 1996; 2006, p. 115).

Must libertarians retreat, then, to the obvious point about indeterministic control that depends on late indeterministic agent-internal processes—that it, unlike deterministic control, is incompatible with the truth of determinism and therefore sufficient for the satisfaction of a necessary condition for free action, if incompatibilism is true? Perhaps not. Perhaps it can be shown that no Frankfurt-style case hits its mark.⁴ But even if some Frankfurt-style cases are successful (in the spheres of moral responsibility and free action), libertarians may try to explain what is attractive about actions that are indeterministically caused by their proximal causes—attractive to believers in free will, that is—beyond the fact that their occurrence is incompatible with the truth of determinism.

The capacity to make and execute intelligent decisions that are indeterministically caused by their proximal causes gives agents who have it a measure of independence from the past (see Mele 1996; 2006 pp. 100–101).⁵ Agents with this capacity can make intelligent contributions to their world of such a kind that it is false that their every thought and action is part of a deterministic causal chain that stretches back for millions of years. They can make intelligent contributions that are not ultimately deterministically caused products of the state of the universe in the distant past—or even of the state of the universe (which includes agents and states of agents) just before the contribution is made or begins. Obviously, the falsity of determinism alone does not bestow this capacity on agents (that is, beings that act). After all, fruit flies do not have this capacity, even if our world is indeterministic. And fruit flies act: they fly, eat, and so on.

Of course, when the topic is whether the falsity of determinism may contribute to free will, comparing human beings to flies is not particularly

helpful. Depending on how amounts of control are to be measured, it may be claimed that having the capacity just mentioned does not give an agent any *more* control than he would have, other things being equal, in a deterministic world. Even if this is true, an agent's having the capacity at issue does open up to him a *kind* of control that he cannot have in a deterministic world (see chapter 7). Now, someone might argue that an agent's having indeterministic control contributes to his having free will only if it contributes to his having *more control* than he would have if he had only deterministic control (or only guidance control). But I have already examined an argument of this kind and found it deficient (chapter 7).

Some philosophers who believe that free will requires agent-causal powers may be thinking that if powers available to event-causal libertarians were sufficient for free will, then incompatibilism would be a hard sell. As such a philosopher might put it, if an indeterministic agent with maximal event-causal libertarian powers has no more control than some agents in deterministic worlds, there is no good reason to believe that the former agent can have free will whereas agents in deterministic worlds cannot. One point to make in reply is that proponents of the more-control argument (see chapter 7) have not presented us with a way of measuring deterministic and indeterministic control on the same scale and have not substantiated the claim that the indeterministic powers at issue provide for no more control than the deterministic ones. A second point, at least as I see it, is that incompatibilism *is* a hard sell; I certainly have not been persuaded by the arguments offered for it (Mele 1995a, 2006).⁶ If someone were to produce a knockdown argument for incompatibilism, we could inspect the argument to see whether it entails that free will requires something that no event-causal libertarian has the resources to offer.

Why would anyone be an event-causal libertarian? Part of the explanation may be that some believers in free will do not regard the argument(s) that persuade them that incompatibilism is true as forcing them all the way to agent causation and share Randolph Clarke's opinion that various arguments collectively "incline the balance against the possibility of . . . agent causation" (2003, p. 209) or his belief that there is no evidence for the existence of agent causation (pp. 206–7) or Derk Pereboom's belief that there is powerful evidence against its existence (2001, chap. 3). These believers in free will may understandably seek a libertarian position with the following two features: its commitments do not include any metaphysical or conceptual impossibilities, and its positive side is supported by evidence. The positive side of event-causal libertarianism is the thesis that there are indeterministic agents who

sometimes act freely. It is what I identified as libertarianism's "pro-free-will" thesis interpreted in light of the event-causal libertarian's appeal to agent-internal indeterminism. The more specific version of this thesis that primarily concerns me here is the one I labeled *LFT*: There are indeterministic agents who sometimes act freely when their actions are not deterministically caused by proximal causes.

Even if worries about the possibility of agent-causal powers were set aside, the issue about evidence would still be very important. Libertarians contend that, in fact, some human beings sometimes act freely. And defending the claim that human beings sometimes engage in a kind of action that requires the existence of a species of causation for which there is no evidence is, to put it mildly, an unpromising project. So it should not be surprising that some incompatibilist believers in free will have stopped short of agent-causal libertarianism. Also, whereas, as far as I can tell, Clarke is right to say that there is no evidence for the existence of agent-causal powers, some scientists claim that they have found evidence of indeterministic behavior-producing brain processes in animals, processes that can be part of the evolutionary heritage of human beings.

Is event-causal libertarianism an unstable pro-free-will position—an uninhabitable halfway house—lying between compatibilism and agent-causal libertarianism? If the more-control argument or the same-control argument were sound (see chapter 7), the view would be unstable in this way. The same is true of Pereboom's disappearing agent objection (see chapter 8). Regarding *LFT* in particular, the point is true as well of Clarke's argument for the falsity of *LDF*—the thesis that even if the difference between what an agent does at *t* in one possible world and what he does at *t* in another possible world with the same past up to *t* and the same laws of nature is just a matter of luck, the agent may perform a directly free action at *t* in both worlds—if, as I believe *LFT* hinges on *LDF* (see chapter 6). But I have shown that these four arguments are unpersuasive. Perhaps some other argument will succeed where these arguments fail. But I have not seen that argument yet.

If it were shown that event-causal libertarianism is unstable in such a way that those who occupy that position should feel serious pressure to move either to compatibilism or to agent-causal libertarianism, what would happen? Given what Clarke and Pereboom tell us about agent causation, agent-causal libertarianism looks like a fast track to the view that no one has free will.⁷ So what would happen would seem to depend a lot on how strong event-causal libertarians think the arguments for incompatibilism are.

A remark by Peter van Inwagen that I mentioned earlier is interesting in this connection. In the final paragraph of a book in which he argues at great length for incompatibilism, van Inwagen writes that “it is conceivable that science will one day present us with compelling reasons for believing in determinism. Then, and only then, I think, should we become compatibilists” (1983, p. 223). Accepting compatibilism is an option for event-causal libertarians who are shown that their view is untenable, provided that they are not entirely convinced that incompatibilism is true.

I am getting ahead of things. No one has shown that event-causal libertarianism is untenable or unstable. And, if there is no evidence for the existence of agent-causal powers, “unstable” is an exceedingly kind assessment of agent-causal libertarianism. Once again, libertarians claim not merely that free action is possible, but also that it is actual; and the contested claim that it is actual should be supported by evidence.

9.3. Event-Causal Libertarianism: Being Positive

I asked what contribution indeterministic agent-internal processes in an action-producing stream might make to free will (if incompatibilism is true) beyond being sufficient for the falsity of determinism. In the preceding section’s discussion of Frankfurt-style cases, I mentioned some independence from the past in this connection. Presumably, we ourselves are rarely in Frankfurt-style situations; and when we are not, the door is open to more than the independence at issue, if there are late indeterministic agent-internal processes in some of our action-producing streams. For example, at a time at which we decided to *A*, things might have been such that in another possible world with the same past up to that time and the same laws of nature, we decide to do something else instead. In such situations, the alternative possibilities available go beyond the constrained kind open in Frankfurt-style cases (for example, deciding on one’s own to steal my car versus deciding to steal my car because one was compelled so to decide).

It may be replied that if any Frankfurt-style cases succeed in the sphere of directly free action, then these less constrained—or *thicker*—alternative possibilities are not required for directly free action and therefore contribute nothing to it. The reply may be taken a step further: since Frankfurt-style cases show that alternative possibilities of the thicker kind are not required for directly free action, and since the more constrained—or *thinner*—kind do not give libertarians what they need in the sphere of alternative possibilities, libertarians should give up their incompatibilism.

This is a lot to digest all at once. One point to notice is that what particular libertarians take themselves to need in the sphere of alternative possibilities should be expected to be bound up with what persuades them that incompatibilism is true. Next on the agenda is an illustration of this point.

Imagine a libertarian who regards a fleshed-out version of *ZAF* below as a conclusive argument for a thesis that any compatibilist about free will would reject and is not persuaded by any other style of argument for that thesis (nor for any incompatibilist thesis that entails it). The point of departure for *ZAF* is a story set in a deterministic world in which a goddess, Diana, creates a zygote, *Z*, in a woman, Mary. Diana “combines *Z*’s atoms as she does because she wants a certain event *E* to occur thirty years later. From her knowledge of the state of the universe just before she creates *Z* and the laws of nature of her deterministic universe, she deduces that a zygote with precisely *Z*’s constitution located in Mary will develop into” an agent, Ernie, who *A*-s thirty years later, thereby bringing about *E* (Mele 2006, p. 188). When Ernie *A*-s, he satisfies an attractive compatibilist set of proposed sufficient conditions for free action and moral responsibility. In a modified version of the story, Diana’s goal in creating Ernie is his performing all the actions she deduced he would perform—that is, every action he ever performs (Mele 2006, p. 190). Her purpose in creating Ernie is to create a being who will perform exactly those actions. This is the operative version here.

Here is *ZAF*.

1. Ernie lacks free will.
2. Concerning free will of the beings into whom the zygotes develop, there is no significant difference between the way Ernie’s zygote comes to exist and the way any normal human zygote comes to exist in a deterministic universe, and no other difference between Ernie and any of the other beings at issue can make it the case that any of them—unlike Ernie—has free will.
3. So in no possible deterministic world in which a human being develops from a normal human zygote does that human being have free will.

A fleshed-out version of *ZAF* would include an argument for premise 2. An argument for that premise might feature the idea that just as Ernie has no say about the conditions in place at the time of his conception and no say about the laws of nature, neither do the other agents at issue, and everything the agents—Ernie and the others—do later is part of the unfolding of their initial conditions. The fleshed-out version might also include the results of efforts to anticipate and rebut various objections to premise 1 (a premise that

is claimed to have intuitive appeal). But it is hard to see why an adequate defense of either premise would need to entail or presuppose the claim that indeterministic alternative possibilities of the thicker kind mentioned above are necessary for free action.

I mentioned “global” Frankfurt-style cases in chapter 5. In cases of this kind, none of the agent’s actions (including decisions) are produced by the potentially compelling processes in the vignettes; but *whenever* the agent acts, such a process would have caused his action if the agent had not, at just the right time, performed an action of the pertinent type on his own. If indeterministic cases of this kind hit their mark in the sphere of directly free action, an agent whose alternative possibilities are never of the thicker kind may perform directly free actions. A libertarian who is persuaded that compatibilism is false by a fleshed-out version of *ZAF* may also believe—consistently—that indeterministic global Frankfurt-style cases do hit their mark in this way. Neither Ernie nor an agent in an indeterministic global Frankfurt-style case—call him Frank—ever has the thicker kind of indeterministic alternative possibility; and even so, the imagined libertarian can consistently claim that Frank sometimes performs directly free actions while denying that Ernie ever does. Imagine that the series of actions Frank performs in his indeterministic universe over his lifetime matches the series Ernie performs in his deterministic universe over Ernie’s lifetime. If Ernie donates \$100 to a fund for orphans at t_1 , so does Frank; if Frank tells a little white lie at t_2 , so does Ernie; and so on. Even so, Frank and Ernie are very different. Ernie is an agent in a deterministic world, and he was created to do everything he does. Frank is an indeterministic agent in an indeterministic world, and the only monkey business in his story is a global failsafe device that never takes over.

As some people conceive of directly free action, it may require something impossible—an agent’s having a kind of control over what he does that is indeterministic and leaves nothing to chance. It may be claimed that when and only when an agent exercises this kind of control is it truly up to him what he does and that agents act directly freely only when it is truly up to them what they do. Why is the kind of control at issue impossible? Because indeterministic control in the absence of chance is impossible. If, for example, an agent exercises direct indeterministic control in deciding to A , and he makes his decision at t , then (setting aside Frankfurt-style cases) in another possible world with the same past up to t and the same laws of nature, he does not decide at t to A : there was, right up to the moment of decision, a chance that he would not decide at t to A .⁸ This is something that any libertarian who holds that directly free actions have proximal causes and cannot be

deterministically caused by their proximal causes needs to learn to live with.⁹ A libertarian who does learn to live with this may (eventually) feel no need to endorse a requirement for directly free action that no one can possibly satisfy. And such a libertarian may consistently continue to believe that free will is incompatible with determinism.¹⁰

I identified the positive side of event-causal libertarianism as the thesis that there are indeterministic agents who sometimes act freely—thesis *LF*, for short. As usual, my primary concern is *directly* free action. Typical libertarians endorse what may be termed a *time-of-action* alternative possibility requirement for directly free actions. They can reasonably disagree about the content of the strongest acceptable such requirement. Here is one candidate: An action *A* performed by an agent *S* at a time *t* in a possible world *W* is directly free only if there is another possible world with the same laws of nature as *W* and the same past up to *t* in which (*O*₁) *S* instead performs some other directly free action at *t*. Libertarians who regard *O*₁ as too demanding have other options. They include, but are not limited to, the following substitutions for *O*₁:

*O*₂. *S* instead performs some other action at *t*.

*O*₃. *S* does not do *A* at *t*.¹¹

*O*₄. Either (1) *S* does not do *A* at *t* or (2) *S* *A*-s at *t*, but not on his own.

These various options are associated with a more ambitious thesis than *LF*—namely *LFT*. Here it is again: There are indeterministic agents who sometimes act freely when their actions are not deterministically caused by proximal causes. As I explained in chapter 7, typical present-day compatibilists have a stake in resisting threats to *LF*, just as event-causal libertarians do, and any compatibilists who believe that some actions are not deterministically caused by their proximal causes and that this does not preclude their being free actions—including, of course, directly free actions—have a stake in resisting threats to *LFT*.

Is there a way to prove that compatibilists and event-causal libertarians should move even closer together—perhaps even completely together? Can it be proved that event-causal libertarians should give up their negative thesis (incompatibilism) or that compatibilists should accept it? It might be said that libertarians who restrict themselves to metaphysical and conceptual possibilities and to claims for which there is some evidence should see that the falsity of determinism makes no contribution to free will. Recall Pereboom's claim that "event-causal libertarianism lacks any significant advantage over

compatibilism in securing moral responsibility” (2001, p. 55), and recall that I am using “free action” as shorthand for moral-responsibility-level (directly) free action. If the following proposition is true, then event-causal libertarians—and everyone else—should become compatibilists: There are indeterministic agents who sometimes perform directly free actions and the falsity of determinism makes no contribution to its being true that these actions are free actions.

A convincing argument for the following proposition would settle the issue about whether event-causal libertarians should abandon their position.

Compatibilism or Bust. Necessarily, either compatibilism is true or free action is impossible.¹²

In the absence of a convincing argument for this bold proposition, the possibility of libertarianism is an epistemically open option. And if no one knows that no human agents are suitably indeterministic agents, then in the absence of a convincing argument for the following proposition, event-causal libertarianism is epistemically open:

Compatibilism Wins a Battle. Necessarily, if there is a true conceptually sufficient condition for an action’s being free that includes the action’s being indeterministically caused, then compatibilism is true.

In my opinion, the strongest theoretical challenge to *LFT* is posed by the problem of present luck. I have replied to that challenge elsewhere (Mele 2006, chap. 5), and I have more to say about it in chapter 10. The strongest empirical challenge to *LF* and *LFT* is to find good evidence of relevant indeterministic processes; and we may have some indirect evidence of that now, if we have direct evidence of indeterministic brain processes in behavior-producing streams of some lower animals.

Another challenge to event-causal libertarianism focuses on the conjunction of its positive and negative sides. Like the theoretical challenge mentioned in the preceding paragraph, this challenge highlights chance or luck. Crudely put, the challenge is to explain how, *given that compatibilism is false*, mixing indeterministic luck or chance into event-causal processes and into the best compatibilist proposals about what suffices for an action’s being free can yield conceptually sufficient conditions for free action. Now, as I have mentioned, I am not an incompatibilist. (I am agnostic about compatibilism.) So the present challenge is not a challenge for me. Of course, even though

(by definition) all libertarians are incompatibilists, no actual libertarian is *essentially* an incompatibilist. If Bob Kane were to reject incompatibilism, he would still be Bob Kane. Event-causal libertarians who do not currently have what they regard as a convincing reply to the challenge at issue may have a variety of options, depending on their philosophical commitments. Options include accepting compatibilism, becoming agnostic about it, becoming agnostic about free action, accepting the view that free action is impossible, and searching for an answer to the challenge. Event-causal libertarians are not boxed in by their incompatibilism unless they are absolutely convinced that it is true.

Recall van Inwagen's assertion that when, and only when, science presents us "with compelling reasons for believing in determinism . . . should we become compatibilists" (1983, p. 223). Given that he is more confident that we sometimes act freely than that compatibilism is false, I suppose he would be willing to modify his assertion if philosophy were to present us with compelling reasons for believing that *Compatibilism or Bust (CoB)* is true. A convincing argument for *CoB*—a major breakthrough, indeed—would presumably move him to endorse compatibilism even if science presents us with compelling reasons for believing that determinism is *false*. Of course, the truth of compatibilism itself is compatible with the truth of *LF* and *LFT* (even if the truth of some detailed compatibilist position is not). And my interest in event-causal libertarianism in this chapter is an interest in its positive side.

Is the operation of all indeterministic action-producing processes of the kinds event-causal libertarians may appeal to in their theories about how free actions are produced actually incompatible with free action? Elsewhere, I have argued that the case for an affirmative answer is not compelling (Mele 2006, chap. 5). Is the operation of these processes incompatible with the truth of the conjunctive proposition that compatibilism is false and there are free actions? As I mentioned, I have not yet seen a convincing argument for an affirmative answer; and for the reasons I adduced in chapter 7, I am not moved by arguments for a *yes* answer that are formulated in terms of amounts of control or that feature the claim that there is no difference between regulative control and guidance control. Is the conjunction composed of incompatibilism and *LF*—or *LFT*—true? I leave that question to libertarians, other incompatibilists, and compatibilists. Having the option of doing so is one of the advantages of being agnostic about compatibilism.

Here it may be useful to repeat a point I made in chapter 7. I mentioned Pereboom's agnosticism about the metaphysical and conceptual impossibility

of agent causation and Clarke's assertion that various arguments collectively "incline the balance against the possibility of . . . agent causation" (2003, p. 209). If agent causation is impossible, then any bar for free will that includes agent causation is impossibly high, and the distance between it and any ordinary compatibilist bar for free will is incalculable. The same is true of event-causal libertarian bars for free will that are achievable in principle, which may help to explain why some philosophers who insist that free will depends on agent causation are unimpressed by differences between indeterministic and deterministic event-causal control. Because both kinds of control are incalculably distant from any impossible species of control, they may appear to be so close to each other that any difference between them cannot have an interesting bearing on free will.

9.4. Conclusion

If incompatibilism is true, what contribution might late indeterministic agent-internal processes in decision-producing streams make to there being directly free decisions beyond being sufficient for the falsity of determinism? The answer I offered is disjunctive. In this concluding section, I summarize it.

Suppose that indeterministic Frankfurt-style cases cannot hit their mark and that, as some libertarians claim, an agent directly freely decides at t to A only if, in another possible world with the same past up to t and the same laws of nature, he performs some alternative action at t . The existence of late indeterministic agent-internal processes in decision-producing streams may support the satisfaction of this alleged necessary condition of directly free decision. And if these processes are such that although they result at t in a decision to A , there are other possible worlds with the same past up to t and the same laws in which they result instead in a decision to B , they ensure the satisfaction of the condition at issue. Furthermore, the mere falsity of determinism does not ensure this even in the case of very similar agents. (So the envisioned contribution the processes at issue make to directly free decisions goes beyond their being sufficient for the falsity of determinism.) Indeterministic worlds are conceivable in which, for any decision to A made at any time t , by t minus 100 milliseconds there was no chance (and no chance*; see note 10) of the agent performing any alternative action at t . Perhaps something else might have happened at t —for example, the agent might be dead or otherwise incapable of acting by t . But then, of course, whatever happened at t was not the agent's performing an alternative action.

Suppose now that some indeterministic Frankfurt-style case does hit its mark and persuades some libertarians to back away from the alleged necessary condition for directly free decision mentioned in the preceding paragraph. Suppose that these libertarians retreat to the following condition: an agent directly freely decides to *A* only if the proximal causes of his decision indeterministically cause it.¹³ If there are late indeterministic agent-internal processes in decision-producing streams, these processes are at work in the production of decisions that are indeterministically caused by their proximal causes. Such decisions satisfy the condition at issue.

I have not said much about how an event-causal libertarian might want to flesh out or interpret the more modest alleged necessary condition at issue. One way to do it appeals to idealized laws of a certain kind. These laws link earlier agent-internal events in decision-producing streams to subsequent events, and they more directly link events near the end of the process to decisions. The laws are idealized in the sense that they assume that there is no interference from outside the stream. An event-causal libertarian may require for directly free decisions that in addition to the agent's having an indeterministic brain in an indeterministic world, the agent's brain is such that the idealized laws at issue—including laws linking proximal causes of decisions to decisions—are probabilistic rather than exceptionless. Conceivably, there are late indeterministic agent-internal processes of an appropriate kind in decision-producing streams. Obviously, in the case of decision-making agents, the falsity of determinism alone does not suffice for the satisfaction of this alleged necessary condition for directly free decision.

The alleged necessary conditions for directly free decision stated in this section are, of course, incompatibilist conditions. They will be rejected by compatibilists. I have no wish to argue here about whether these conditions are true or false. My primary concern has been with what I called the positive side of event-causal libertarianism—that is, *LF* and, more specifically, *LFT*. And, as I explained in chapter 7 and reminded readers once already in the present chapter, typical present-day compatibilists who believe that we are indeterministic agents have a stake in resisting threats to *LF* and compatibilists who believe that some actions are not deterministically caused by their proximal causes and that this does not preclude their being free actions have a stake in doing the same for *LFT*.

Elsewhere, I have developed and explored what I call *soft* libertarian views (Mele 1996, 2006). Soft libertarianism is the thesis that “free action and moral responsibility [may be] compatible with determinism but . . . the falsity of determinism is required for . . . more desirable species of” these things (Mele

2006, p. 95). Some philosophers have speculated about why I have done this (Nelkin 2007). Part of the answer is that I would like to understand what various positive libertarian ideas might have going for them without taking a stand on the negative side of libertarianism (incompatibilism). This chapter was written in the same spirit.¹⁴

Notes

1. Not everyone who has written about agent causation takes it to depend on the falsity of determinism. See Markosian 1999; Nelkin 2011; and Pereboom 2015.
2. Reminder: libertarians who reject the idea that there are indirectly free actions should ignore the word “directly” in “directly free.”
3. *PAP* reads as follows: “A person is morally responsible for what he has done only if he could have done otherwise” (Frankfurt 1969, p. 829). For disambiguation, see Mele 2006, chap. 4.
4. For the record, I have argued (and continue to believe) that what are sometimes called Mele-Robb Frankfurt-style cases (Mele and Robb 1998) undermine a variety of alternative-possibility principles both in the sphere of moral responsibility and in the sphere of free action (Mele 2006, chap. 4). For replies to various objections to Mele-Robb cases, see Mele and Robb 2003.
5. The “hard libertarian” view presented in Mele 1996 encompasses what was subsequently called “source incompatibilism” (McKenna 2001). I should add that early indeterministic agent-internal processes can also contribute to an agent’s having some independence from the past.
6. The potential buyers I have in mind are philosophical experts on free will who are either compatibilists or agnostics about compatibilism. There is no point in trying to sell incompatibilism to experts who already endorse it.
7. As I mentioned, Pereboom’s argument that free will depends on agent causation is part of his argument for the thesis that free will does not exist (2001, 2014). There is no argument in Clarke 2003 for the nonexistence of free will; Clarke’s concern there is to assess the conceptual adequacy of libertarian views.
8. In a Frankfurt-style version of the example, if the agent decided on his own at t to A , there was a chance that he would decide at t to A without making that decision on his own.
9. It may be claimed that an agent can have direct indeterministic control over what he intends even though there is no chance of his deciding to A without intending to A . As I understand deciding to A (chapter 2), it is conceptually sufficient for intending to A . Accordingly, in my view, there is no chance of the following: although S decides at t to A , he does not intend at t to A . But (setting aside Frankfurt-style cases and nonactionally acquired intentions) before t , there was a chance that S would not at t have an intention to A , if he formed that intention at t in deciding

to *A* and exercised direct indeterministic control in so deciding; for, right up to *t*, there was a chance that he would not decide then to *A*. (Again, in a Frankfurt-style version of the case, if *S* decided on his own at *t* to *A*, there was a chance that he would decide at *t* to *A* without so deciding on his own.)

10. Is there a loophole? Might it be that although a normal human agent exercised direct indeterministic control at *t* in deciding to *A*, there was no antecedent objective probability that he would decide at *t* to *A* and no such probability that he would do any of the other things he might instead have done at *t*? If so, and if chances are identified with objective probabilities, then there was no chance that he would decide at *t* to *A* (even though he did so decide) and no chance that he would do any of these other things. (Presumably, there was an objective probability of 0 that he would jump over the moon at *t*; but, by hypothesis, of the things he might have done at *t*, there was no objective probability that he would do them.) If cases of the sort at issue are possible, and if chances are to be identified with objective probabilities, my claim that “there was, right up to the moment of decision, a chance that he would not decide at *t* to *A*” can be modified to accommodate them: replace “chance” with “chance*” and define the latter disjunctively as chance or the non-chance openness allegedly present in these cases. (Randy Clarke raised the issue discussed in this note.)
11. Obviously, that *S* does not do *A* at *t* does not entail that *S* does nothing at all at *t*. In typical cases of the kind under consideration, an agent who satisfies *O*₃ also satisfies *O*₂.
12. The first disjunct obviously does not make free action depend on determinism; that disjunct leaves it open that there are free actions in indeterministic worlds. In terms of possible worlds, what it asserts is that there are possible worlds in which determinism is true and agents sometimes act freely.
13. Other moves are possible. In Frankfurt-style cases, agents *A* on their own and they could have done otherwise than *A* on their own. It might be claimed that the agent freely *A*-s on his own partly because he could have done otherwise than *A* on his own (for a related suggestion about moral responsibility, see Naylor 1984; Robinson 2012; and van Inwagen 1983, p. 181).
14. For comments on a draft of Mele 2013c, from which this chapter derives, I am grateful to Randy Clarke and Stephen Kearns.

Two Libertarian Theories

OR WHY EVENT-CAUSAL LIBERTARIANS SHOULD
PREFER MY DARING LIBERTARIAN VIEW
TO ROBERT KANE'S VIEW

SOME LIBERTARIAN VIEWS feature agent causation, others maintain that free actions are uncaused, and yet others—event-causal libertarian views—reject all views of these two kinds and appeal to indeterministic causation by events and states. This chapter explores the relative merits of two different views of this third kind. One is Robert Kane's prominent view (1999b, 2011, 2014), and the other is the “daring libertarian” view that I floated in *Free Will and Luck* (Mele 2006). (I labeled the view “daring” to distinguish it from a “modest” libertarian view that I floated in Mele 1995a.) I say “floated” because I have never been a libertarian. I do not endorse incompatibilism; instead, I am agnostic about it. But if I were a libertarian, I would embrace my daring libertarian view (or *DLV*, for short). This chapter's thesis is that event-causal libertarians should prefer *DLV* to Kane's “dual or multiple efforts” view (2014, p. 209).¹ As I have mentioned, I argue in Mele 2006 that if there is a viable libertarian view, it is of the event-causal kind. So I do not regard my focusing on event-causal libertarianism here as problematic.

10.1. Kane's Concurrent-Efforts View

Kane distinguishes among “three freedoms” (2008, p. 142). He asserts: “Free acts may be”

- (1) acts done voluntarily, on purpose and for reasons that are not coerced, compelled or otherwise constrained or subject to control by other agents.

- (2) acts [free in sense 1 that are also] done “of our own free will” in the sense of a will that we are ultimately responsible (UR) for forming.
- (3) “self-forming” acts (SFAs) or “will-setting” acts by which we form the will from which we act in sense 2. (2008, p. 143)²

Kane observes that free actions of type 1, as he conceives of them, are compatible with determinism and that free actions of types 2 and 3 are not (p. 143). All free actions of type 3, as Kane conceives of them, are indeterministically caused by their proximal causes, and only agents who perform free actions of type 3 can perform free actions of type 2.

In chapter 6, I used the label “basically free action” for any free actions—free *A*-ings—that occur at times at which the past (up to those times) and the laws of nature are consistent with the agent’s not *A*-ing then. My focus in this chapter is on basically free actions that are indeterministically caused by their proximal causes. If there are such actions, it is possible that all actual actions of this kind are self-forming actions. But whether an action is self-forming or not depends on its effects on the agent’s “will”; and if basically free actions are possible, we can imagine basically free actions that immediately precede the agent’s death and therefore have no effect on the agent’s will.

An agent performs a basically free action *A* at a time *t* only if there is another possible world with the same past up to *t* and the same laws of nature in which he does not do *A* at *t*.³ In some cases, the *A*-ing is an action of deciding (or choosing) to do something or other—deciding to cheat right then, for example. And in many such cases, there is another possible world with the same past up to *t* and the same laws of nature in which, at *t*, the agent decides to do something else instead. Reflection on such pairs of worlds has led some philosophers to worry that what the agent decides to do is too much a matter of luck for the agent to be morally responsible for the decision and to have made it freely. In my own formulation of the worry, as I have explained, the cross-world difference in decisions at *t* is just a matter of luck, and, for example, the agent’s deciding at *t* to cheat is partly a matter of luck (Mele 2006, pp. 8–9, 54–55, 114, 132–33).

I have never claimed that the luck here is incompatible with basically free action or with the agent’s being morally responsible for the action. Instead, I have offered a solution to the worry that acknowledges the presence of luck at the time of action (Mele 2006, chap. 5). Kane offers another solution (1999b, 2011, 2014). His proposed solution also acknowledges the presence of luck: “The core meaning of ‘He got lucky,’ which *is* implied by indeterminism, I suggest, is that ‘He succeeded *despite the probability or chance of failure*’; and

this core meaning does not imply lack of responsibility, if he succeeds" (Kane 1999b, p. 233).

Elsewhere, I have explained how Kane's libertarian view evolved over the years in response to some worries about luck (Mele 2006, pp. 75–76). Here I cut to the chase more expeditiously. Kane's proposed solution to one such worry, in cases in which the actions at issue are decisions, features the idea that the agent simultaneously tries to make each of two or more competing choices or decisions (1999b).⁴ In this chapter, to keep things relatively simple, I restrict attention to cases in which only two competing choices are in the running. Regarding such cases, Kane claims that because the agent is trying to make each choice, she is morally responsible for whichever of the two choices she makes and makes it freely (pp. 231–40), provided that "she endorse[s] the outcome as something she was trying and wanting to do all along" (p. 233). If Kane is right, he has provided a successful answer to a certain challenge about luck (see Mele 2006, chap. 3; see also chapter 6 above)—at least in scenarios of a certain kind.

Part of the inspiration for Kane's position is the observation that "indeterminism [sometimes] functions as an obstacle to success without precluding responsibility" and free action (1999b, p. 227). In one of his illustrations, "an assassin who is trying to kill the prime minister . . . might miss because" his indeterministic motor control system leaves open the possibility that he will fire a wild shot. Suppose the assassin succeeds. Then, Kane says, he "was responsible" for the killing "because he intentionally and voluntarily succeeded in doing what he was *trying* to do—kill the prime minister" (p. 227). It may be claimed, similarly, that the indeterminism in the scenario does not preclude the killing's being a free action. If these claims are true, they are true even if the difference between the actual world at a time during the firing and any wild-shot world that does not diverge from the actual world before that time is just a matter of luck.

Kane contends that "libertarian views in general must try to show that whatever chance may be involved in undetermined choices need not undermine free agency and responsibility" (2014, pp. 207–8). He also contends that to show this one must go beyond my daring libertarian view (to be described in sections 10.3 and 10.4) and defend the claim that "the agent *makes one set of reasons win out* over the other at the moment of choice, so that the agent can be fully responsible for causing it to be the case that one choice rather than the other is made, despite the indeterminism" (p. 208). Here he appeals to his *concurrent-efforts* idea: "the agent *makes one set of reasons prevail* over the other *by* making an effort to do so against the competing effort to make a

contrary choice” (p. 208). The assassin succeeds at something he is trying to do, despite the luck involved. Similarly, Kane says, the person who makes dual efforts to choose may succeed—whichever of the two competing choices he makes—at doing something he is trying to do, despite the luck involved. And his succeeding includes his making one set of reasons prevail over the other.

10.2. Probing Kane’s View

As I pointed out elsewhere, although Kane is ordinarily pretty sensitive to the phenomenology of agency, trying to choose to *A* while also trying to choose to do something else instead seems remote from ordinary experience (Mele 2014a, p. 43). We may occasionally have an experience of trying to bring it about that we choose to *A*. (For example, someone who knows that it would be best to quit smoking but who has not yet chosen to quit may try to vividly represent to himself the most important reasons for quitting, including the dangers of not quitting, with a view to bringing it about that he chooses to quit.) But how many of us have experienced simultaneously trying to bring it about that we choose a particular course of action and trying to bring it about that we choose a competing course of action instead? If such dual efforts never occur, they never underwrite free choices.

Obviously, someone who grants that we never experience dual efforts to choose may wish to posit such efforts anyway in an attempt to solve a theoretical problem (see Kane 2011, pp. 391–92; 2014, pp. 193–202, 208–9). But alternative routes to a solution should also be explored. If concurrent efforts of the kind at issue never happen, is an event-causal libertarian unable to explain why some choices that are indeterministically caused by their proximal causes are free? I take up this question in subsequent sections.

In response to the phenomenological worry, Kane asserts that “introspective evidence cannot give us the whole story about free will” (2014, p. 197). But, of course, no one who expresses the worry claims that such evidence can give us the whole story. Instead, people like me have their doubts about whether normal human beings ever simultaneously try to make each of two (or more) competing choices or decisions, and phenomenological considerations are among the considerations we cite in support of the doubts.

I turn from this empirical issue to a theoretical question. Even if people sometimes do make dual efforts of the kind at issue, would that give Kane the result he wants? Consideration of some cases featuring dual efforts of some other kinds will prove useful.

Here is a warm-up case. Ann knows that either of a pair of targets will suddenly disintegrate before a bullet fired at it hits it and that it is undetermined which target will do this. She is promised a prize of \$10 for hitting target 1 and \$20 for hitting target 2. Ann is ambidextrous and an expert with firearms. She fires simultaneously at each target, shooting at target 1 with the pistol in her right hand and at target 2 with the pistol in her left hand. As luck would have it, target 2 disintegrates and Ann hits target 1.

This story differs from stories about Kane-style dual efforts to choose in three potentially noteworthy ways. It is not a story about alternative choices, the agent's efforts end before success is achieved (Ann's efforts end when she pulls the triggers), and the efforts do not hinder one another.

Here is a story in which the second difference is eliminated. Beth is promised a prize of \$10 for fully depressing the V key on a computer keyboard with her left index finger and \$20 for fully depressing the M key on the same keyboard with her right index finger. She knows that one key or the other will stick (and so will not fully depress) and that it is undetermined which. She also knows that if she opts to press both keys she must press them simultaneously in order not to be disqualified. Her plan is to press each key simultaneously. She tries to press the V key all the way down with her left index finger while also trying to press the M key all the way down with her right index finger. As luck would have it, the V key sticks and Beth fully depresses the M key.

The next story eliminates two of the three differences. Cathy's situation is like Beth's except that her index fingers are linked together by a fancy collection of plastic strings and pulleys. Moving either finger downward makes it harder to move the other finger downward. Her fingers are an inch above the keys, and her plan is to try to press each key at the same time. As luck would have it, the M key sticks and she fully depresses the V key.

Here is an obvious point about these cases. It is not up to the agent which target she hits or which button she fully depresses. Is it up to agents what they choose in Kane-style cases of dual efforts to choose? Or is what happens pretty well understood on the model of my third case? In that case, we have two simultaneous attempts, each of which hinders the other, and one of them happens to succeed while the other one happens to fail.

Recall Kane's claim that "the *agent makes* one set of reasons prevail over the other *by* making an effort to do so against the competing effort to make a contrary choice" (2014, p. 208). One might try to do justice to this claim by representing the agent as trying to make reasons R₁ prevail over reasons R₂ while also trying to make reasons R₂ prevail over reasons R₁. But this does not capture the idea that the agent is making the former effort "against" the

latter effort, if what is being claimed is that the agent represents the former effort as being undertaken against the latter.⁵ If that is what is being claimed, we can say the following: the agent is trying to make reasons R_1 prevail over reasons R_2 and to prevent his contrary effort from making R_2 prevail over R_1 while also trying to make reasons R_2 prevail over reasons R_1 and to prevent his contrary effort from making R_1 prevail over R_2 .

My third case can be modified accordingly. This time, Cathy performs her task while in a brain scanner. She is told that one thing she must do to win either prize is to represent what she is trying to do in a certain way and that the scanner will reveal how she represents her attempts. The relevant part of her instructions read as follows:

One thing you must do to win either prize is to represent yourself as trying to make your reason for fully depressing the V key (that you will get \$10 for doing so) prevail over your reason for fully depressing the M key (that you will get \$20 for doing so) and to prevent your contrary effort from making the latter reason prevail over the former while also representing yourself as trying to make your reason for fully depressing the M key prevail over your reason for fully depressing the V key and to prevent your contrary effort from making the latter reason prevail over the former.

After thinking about this for a while and then trying to represent her options to herself in terms of making one reason prevail over another, Cathy reports that she is almost ready. After a bit more thinking—specifically about the idea of her trying to prevent an effort of hers from being successful—Cathy reports that she is ready. According to the scanner, she represents what she is up to in the specified way and, as luck would have it, the M key sticks and Cathy fully depresses the V key.

In this case, as in the earlier story about her, it is not up to Cathy which button she fully depresses. Which button she fully depresses depends on which button gets stuck, and she has no say at all about which button gets stuck.

Here, Kane may say, we have hit on an important difference between Cathy's story and a case of Kane-style dual efforts to choose. In Cathy's story, the outcome hinges on an external event over which she has no control; but in a Kane-style story, all the work is done by the agent's own activity—his dual efforts. How much mileage can one get out of this difference?

Reflection on my story about Bob and the coin in chapter 6 will help answer this question. Imagine, if you can, that Bob is trying to choose to toss

the coin at noon, as promised, while also trying to choose to cheat and that these efforts are or include efforts to make pertinent reasons prevail over other pertinent reasons and are made “against” each other (Kane 2014, p. 208). In possible world W_1 , Bob’s attempt to choose to cheat succeeds at t . But there is another possible world, W_2 , with the same past up to t and the same laws of nature in which, at t , Bob’s attempt to choose to toss the coin on time succeeds. We get two very different outcomes with no antecedent difference at all.

In W_1 , what is the status at t of Bob’s attempt to choose to toss the coin on time? One possibility is that it has not yet stopped. It is still underway at t , but it is not successful. Another possibility is that this attempt stopped just then. We can say, if we like, that Bob’s attempt to choose to cheat rendered this competing, persisting attempt ineffective in the former scenario without stopping it and that it stopped the competing attempt at t in the latter scenario. (Recall Kane’s claim that “the agent makes one set of reasons prevail over the other by making an effort to do so *against the competing effort* to make a contrary choice” [2014, p. 208; italics altered].) But we should not lose sight of the point that what is going on is such that, in another possible world with the same laws of nature and the same past right up to t , exactly the opposite happens. There is no difference in the efforts—or anything else, for that matter—before t ; and, even so, in W_2 Bob’s attempt to choose to toss the coin on time succeeds (and, if you like, stops the competing attempt at t or renders it ineffective at t without stopping it then). The difference at t between W_1 and W_2 seems to be just a matter of luck. And readers would understandably have doubts about claims that the following things were up to Bob: which of his two efforts to choose succeeded; which set of reasons prevailed; what he wound up choosing when his dual efforts to choose ran their course.

As I see it, to assert that the difference at issue is just a matter of luck is not to assert that Bob’s decision is not a basically free action or not something for which he is morally responsible. In fact, as readers of chapter 6 know, one plank in my response to the problem about luck at issue is the following thesis: (LD) Even if the difference between what an agent decides at t in one possible world and what he decides at t in another possible world with the same past up to t and the same laws of nature is just a matter of luck, the agent may make a basically free decision at t in both worlds.⁶ In Mele 2006, as I have mentioned, I present what I called the problem of “present luck” (p. 66) in the same spirit that someone who hopes for an adequate explanation of why a perfect God would allow all the pain and suffering that exists may vividly present the problem of evil, and I do it without formulating any argument for

a conclusion that is incompatible with *LD*. Part of what I asked for, in effect, was a plausible explanation of the truth of *LD*. Kane has offered the answer that I have been discussing. I will get to the alternative answer that I offered pretty soon.

LD resembles the following thesis about actions in general: (*LG*) Even if the difference between what an agent does at *t* in one possible world and what he does at *t* in another possible world with the same past up to *t* and the same laws of nature is just a matter of luck, the agent may perform a basically free action at *t* in both worlds. Recall my stories about Ann, Beth, and Cathy. It is not up to Ann whether she hits target 1 or target 2, and it is not up to Beth and Cathy whether they fully depress the V key or the M key. Even so, readers who believe that basically free actions are possible may well regard the actions at issue—Ann's hitting target 1, Beth's fully depressing the M key, and Cathy's fully depressing the V key—as basically free. And, of course, if these actions are basically free, the same goes for Ann's hitting target 2, Beth's fully depressing the V key, and Cathy's fully depressing the M key in worlds with the same laws and past in which that is what happens.

There are versions of my stories about Ann, Beth, and Cathy in which their actions have moral significance. For example, we can imagine a version of the keyboard stories in which key presses are means of administering painful shocks to kittens. Fully depressing the M key administers a shock to an adorable gray kitten, and fully depressing the V key does the same to an equally adorable white kitten. Beth and Cathy press the M key in an attempt to shock the gray kitten, and they press the V key in an attempt to shock the white kitten. If the kitten-shocking actions are performed freely, might the agents be morally responsible for them? In Kane's view, as I mentioned, the assassin with an indeterministic motor-control system is morally responsible for killing the prime minister, something he was voluntarily trying to do. Beth and Cathy also voluntarily try to do what they succeed in doing (but while also trying to shock the other kitten). So when Beth and Cathy shock the gray (or white) kitten, Kane may be happy to say that they are morally responsible for doing that. But what he should not—and presumably would not—say is, for example, that Cathy is morally responsible for the fact that she shocked the gray kitten *rather than* the white kitten. She lacks moral responsibility for this contrastive fact.⁷ Her being morally responsible for that fact would require that she is morally responsible both for the fact that she shocked the gray kitten and for the fact that she did not shock the white kitten. And although Cathy's attempt to shock the white kitten failed, she is not

morally responsible for its failing nor for her not shocking this kitten. I return to this matter in section 10.5.

10.3. Comparing the Competing Views

Libertarianism has a negative and a positive side. The negative side is incompatibilism, and the positive side is the claim that at least some human beings sometimes act freely—or act in a way that depends on their having free will. The positive claim is not merely that free will is possible; it is that free will is actual and actually exercised by real people. Now, any typical libertarian view makes it a necessary condition for a directly free action that there is no time at which it is determined (in the “deterministic causation” sense of determined) that the action will occur.⁸ And any typical *event-causal* libertarian view holds that although directly free actions are caused, they are not deterministically caused by their proximal causes. Given the positive side of libertarianism, we have here a commitment about actual human beings: namely, that at least some of us sometimes perform actions that are not deterministically caused by their proximal causes. Because, in any case of action, at least some of the proximal causes of actions are internal to agents, the commitment here for typical event-causal libertarians is to there being indeterminism in some action-producing causal streams right up to the time of action (including choice or decision).

Kane certainly seems to hold that dual (or multiple) efforts of the kind he posits are *required* for directly free choices. This adds a second commitment about actual human beings—at least, those who make directly free choices. This commitment is not merely to what is possible. It is a commitment about what some actual human beings actually do. Kane mentions evidence about dual processing in perception (2014, p. 197); but, to the best of my knowledge, there is no direct evidence of Kane-style dual attempts to choose in normal human beings.

Two things have kept me from being a libertarian. First, I have not been persuaded by any argument for incompatibilism. Second, for reasons I have set out elsewhere (Mele 2006), I take (a naturalistic) event-causal libertarianism to be the most promising brand of libertarianism, and I do not know of strong evidence that human brains work as they would need to work if a theoretically attractive event-causal libertarian view is true. Both of these things stand in the way of my endorsing the daring libertarian view that I floated. But, as I observed, all I want to argue here is that event-causal libertarians should prefer that view (*DLV*) to Kane’s view.

DLV (Mele 2006) is similar to Kane's view. The main difference is that where Kane postulates concurrent competing indeterministic efforts to choose, I postulate an indeterministic effort to decide (or choose) what to do. That effort can result in different decisions, holding the past and the laws of nature fixed. For example, in Bob's story, as I tell it in Mele 2006 and in chapter 6 above, there are no concurrent competing efforts to choose. Instead, there is a possible world in which Bob's effort to decide what to do about the coin toss issues at t in a decision to cheat, and in another world with the same past up to t and the same laws of nature, that effort issues at t in a decision to toss the coin right then. Bob has competing reasons at the time, and the decision he makes—whether it is to cheat or to do the right thing—is made for the reasons that favor it. The cross-world difference at t in what Bob decides seems to be a matter of luck. But it does not seem to be any *more* a matter of luck than a cross-world difference that I identified in a version of Bob's story in which he is trying to choose to cheat while also trying to choose to do the right thing—namely, the difference between the former effort succeeding and the latter effort succeeding.

Kane contends that choices of the sort at issue—choices that issue from one member of a pair (or group) of competing efforts to choose—are 'up to the agent' in the strong sense that the agents have plural voluntary control over whether or not they are made" (2014, p. 202). He comments on the nature of plural voluntary control earlier in the same article:

We are interested in whether [agents] could have acted voluntarily and intentionally in *more than one way*, rather than in only one way, and in other ways merely by accident or mistake. I call such conditions of more than one way voluntariness and intentionality *plurality conditions* for free will, and the power to act in accordance with them *plural voluntary control*. (2014, p. 185)

As Kane grants, when an agent satisfies the conditions set out in *DLV*, "we can . . . say that the choice that results is made by the agent; and we can even say it is voluntary (since uncoerced) and intentional (since knowingly and purposefully made)" (p. 207). Moreover, we can correctly say these things both about the choice the agent makes at t in the actual world and about a competing choice he makes at t in another possible world with the same laws of nature and the same past up to t .⁹ And from this we can infer that *DLV* accommodates "plural voluntary control" over choice-making. If, as Kane says, plural voluntary control in this connection is sufficient for the choice

the agent makes to be up to the agent, then *DLV* also accommodates it sometimes being up to agents what they choose in scenarios of the sort at issue. The point is that, given Kane's own account of plural voluntary control and his own claim about what is sufficient for an agent's choice to be up to the agent, an agent whose making of a particular choice fits *DLV* makes a choice that it was up to him to make.

Even so, Kane finds fault with *DLV*. He writes: "The agent will indeed make one choice or the other at *t*, but which choice the agent makes depends on which reasons 'win out'; and this is undetermined. That the agent decides to do *A* at *t* in one world and *B* in another seems therefore to be a matter of luck or chance" (2014, p. 207).

A pair of observations are in order. The first is about victorious reasons. Call Bob's reasons for cheating *RC* and his reasons for tossing the coin at noon *RT*. In the present context, what it is for *RC* to win out is for Bob to choose for those reasons; and if that happens, then, of course, Bob chooses to cheat. Now, in order for Bob's choice to cheat to be basically free—according to Kane's view and *DLV*—there must be no time at which it is determined that he will choose to cheat and so no time at which it is determined that he will choose for *RC*.¹⁰ So Kane cannot, given his own view, treat the point that which reasons will win out is undetermined as incompatible with Bob's making a basically free choice to cheat. The upshot, of course, is that he cannot consistently claim that this point falsifies *DLV*. If the point falsifies *DLV*, it falsifies Kane's view too.

My second observation is predictable, given some remarks I have already made. When Bob's choice-making occurs in a way that fits *DLV*, the cross-world difference in what he chooses is, in Kane's words, "a matter of luck or chance" (2014, p. 207). But when Bob's choice-making occurs in a way that fits Kane's concurrent-efforts view, the cross-world difference in which of his efforts to choose wins out is no less a matter of luck or chance. Picture Bob's dual efforts to choose as (mutually interfering) processes aimed at targets or goals.¹¹ The target at which his effort to choose to cheat aims (*T₁*) is his choosing to cheat, and the target at which his effort to choose to do the right thing aims (*T₂*) is his choosing to do the right thing. In *W₁* the former process wins out at *t*: *T₁* is hit then and *T₂* is not. And in *W₂*, which has the same past all the way up to *t* and the same laws of nature, the latter process wins out: *T₂* is hit then and *T₁* is not. This difference is no less a matter of luck than the featured cross-world difference when Bob's choice-making accords with *DLV*. More is going on in Kane's vision of things than in mine. He represents the agent as trying to make each of two different competing choices, and I represent him

as simply trying to decide what to do. But the more that is going on in Kane's concurrent-efforts picture does not yield less cross-world luck.

I will say more about *DLV* shortly, and I will compare the main costs of the two views at issue. But I would like to linger for a while over Kane's idea that, in a dual-efforts scenario, the agent makes some reasons prevail over others.

In my story in which Cathy is in a brain scanner, readers may have a hard time imagining her trying to make her \$10 reason prevail over her \$20 reason. Such an attempt would be perverse, and it seems that all she is really up to is simultaneously trying to fully depress the \$10 key and trying to fully depress the \$20 key while also attempting to represent what she is doing in a certain reason-featuring way. So consider the following case. Donna replaces Cathy in the scanner, each full key press is worth \$15, and the two keys assign money to two different charities—one for stray dogs and the other for stray cats. The representation instructions that Donna receives reflect this fact, of course. Further details: Donna is very fond of cats and dogs and very interested in helping them, and she believes that the two charities are equally proficient at achieving their aims.

Donna attempts to follow her instructions about representations. The strategy that she endeavors to implement includes her vividly imagining the plight of an adorable stray kitten in an attempt to make her reasons for helping stray cats prevail and vividly imaging the plight of an equally adorable stray puppy in an attempt to make her reasons for helping stray dogs prevail. Her former attempt also includes rehearsing reasons for helping stray cats and the latter includes rehearsing reasons for helping stray dogs. Donna simultaneously presses both keys. In the actual world, she fully depresses the cat key at *t*, and in another possible world with the same laws and the same past up to the sticking point, she fully depresses the dog key at *t*. We can say, if we like, that in fully depressing the cat key, Donna made her reasons to help cats prevail. And this claim can be counted as true, if we do not read too much into it. But the truth of the claim is utterly compatible with the difference in the two worlds at the time at issue being just a matter of luck. And, in a Kane-style dual-efforts scenario, the same is true of the difference between Bob's making his reasons to cheat prevail in choosing at *t* to cheat and his making his reasons to do the right thing prevail in choosing at *t* to toss the coin. Bear in mind that Bob does not make either set of reasons prevail *before* he makes his choice. Which reasons prevail is up for grabs until he makes his choice, and the prevailing of a collection of reasons is precisely a matter of Bob's choosing for those reasons—that is, his choosing for reasons *RC* to cheat or his choosing for reasons *RT* to do the right thing. Again, in one world one set of reasons prevails at *t*, and in another world a competing set of reasons prevails

at t —and there is no cross-world difference in the reasons, Bob’s efforts, or anything else before t .¹²

As I observed, Kane contends that his view secures the making of choices that “are ‘up to the agent’ in the strong sense that the agents have plural voluntary control over whether or not they are made” (2014, p. 202). And, as I explained, my own *DLV* fares no less well in this regard: agents who fit the latter view can satisfy Kane’s sufficient conditions for plural voluntary control over what they choose. Kane can reply that exercising plural voluntary control over what one chooses is not, after all, sufficient for acting freely because an additional requirement for choosing freely is that the agent was trying specifically to choose what he chose.¹³ But what contribution does a choice’s satisfying this requirement make to its being a free choice? Kane might, by way of analogy, point to the contribution that the assassin’s trying to kill the prime minister (in an example mentioned earlier) makes to the killing’s being a free action. However, the contribution here seemingly consists in the support the occurrence of the trying offers for the claim that the assassin *intentionally* killed his victim, presumably as a means to an end (Kane 1999b, p. 227); and it is commonly recognized that choosing to A is *essentially* intentional (McCann 1986a; Mele 1997). Agents do not need to try to choose to A nor to try to bring it about that they choose to A in order to choose to A ; and the nature of choosing is such that, whenever they choose to A , they intentionally do so: there are no nonintentional choosings to A . If trying to choose to A is supposed to make a contribution to freely choosing to A that goes beyond its contribution to exercising plural voluntary control over what one chooses, Kane has not said what that contribution is.

10.4. The Two Views: Relative Costs

I have made two main points that bear directly on this chapter’s thesis.

1. *On cross-world luck.* When an agent’s choice-making occurs in a way that fits my *DLV*, the cross-world difference in what he chooses is, in Kane’s words, “a matter of luck or chance” (2014, p. 207). When the agent’s choice-making occurs in a way that fits Kane’s concurrent-efforts view, the cross-world difference in which of his efforts to choose wins out is no less a matter of luck or chance.
2. *On empirical burdens.* Kane’s concurrent-efforts view requires more for basically free decisions than *DLV* does. It requires that the agent simultaneously makes competing efforts to choose. Ordinary experience supports

the claim that normal human agents sometimes make an effort to decide what to do. The same cannot plausibly be said for the claim that agents sometimes make concurrent efforts to choose of the kind featured in Kane's view. And, to the best of my knowledge, there is no direct evidence of any kind that normal agents ever make Kane-style concurrent efforts to choose.

The conjunction of 1 and 2 is motivation for this chapter's thesis—that event-causal libertarians should prefer *DLV* to Kane's concurrent-efforts view. Kane's view has a significantly heavier burden—and therefore carries a significantly higher cost—on the empirical front, and, as far as I can see, it has no advantage over *DLV* on the issue of cross-world luck at the time of choice or decision.

Regarding point 1, some people may worry that *DLV* leaves more room for wildly improbable choices than Kane's view does. But the worry is unfounded. Presumably, a typical Kane-style agent's character, values, beliefs, learning history, and the like constrain the options that it is open to him to choose, and there is nothing to prevent these factors from being equally constraining in an agent who fits *DLV*.

Regarding point 2, I have heard it said that Kane is not seriously claiming that agents make simultaneous competing efforts to choose—that all he really has in mind is that the agent has reasons or motives for two or more competing choices and chooses for reasons. Quotations from Kane's work that I have provided here should make it clear that this interpretation is far off the mark, as should a little reflection on the use to which Kane puts his assassin analogy (discussed earlier; see also Kane's table-shattering example: 1999b, p. 227; 2014, pp. 194, 200). Moreover, on *DLV*, agents have reasons or motives for two or more competing choices and choose for reasons in cases of the kind at issue; and Kane makes it clear that he is not willing to settle for this, as I have observed.

10.5. Daring Libertarianism and Agents' Histories

Another difference between *DLV* and Kane's concurrent-efforts view merits mention. Kane tries to find a solution to a worry about present luck in what is happening at and around the time of a decision. This is something he has in common with agent causationists. According to *DLV*, not everything that is needed can be found there; we should also look back in time. A detailed discussion of what a daring libertarian hopes to find in agents' histories is

beyond the scope of the present book. My *DLV* finds in reflection on agents' pasts a partial basis for an error theory about why *some* people may view cross-world luck at the time of decision as incompatible with deciding freely (2006, pp. 111–34). Brief attention to this issue is in order.

My error theory is for a limited audience—people who are attracted to libertarianism and reject agent-causal and noncausal libertarianism. When some such people reflect on stories like that of Bob and the coin, they may ignore the sources of the antecedent probabilities of Bob's choosing to cheat and his choosing to do the right thing. If it is imagined that these probabilities come out of the blue, Bob may seem to be adrift in a wave of probabilities that were imposed on him, and, accordingly, he may seem not to have sufficient control over what he chooses to be morally responsible for his choices. But, as I have explained elsewhere, it is a mistake to assume that "indeterministic agents' probabilities of action are externally imposed" or that such agents "are related to their present probabilities of action roughly as dice are related to present probabilities about how they will land if tossed" (2006, pp. 124–25). If it is known that Bob's pertinent probabilities shortly before noon are shaped by past intentional, uncompelled behavior of his, one may take a less dim view of Bob's prospects for being morally responsible for the choice he makes and his prospects for making it freely.

This is a long story that carries us all the way back to candidates for young agents' earliest basically free actions (see Mele 2006, pp. 111–34). I drew upon it in section 8.5 of chapter 8. I cannot do justice to it here without rehashing a lot of material from Mele 2006. But I will say a bit more about it.

In the course of sketching the problem of present luck in chapter 6, I used an analogy with a genuinely random number generator. It is natural to want to respond by pointing out differences between mindless number generators and indeterministic human decision-makers. Here is one difference: whereas what random number a genuinely random number generator generates next is causally independent of its earlier productions of random numbers, our decisions often seem to be causally influenced by earlier decisions we have made (and such influence seems not to depend conceptually on our being deterministic decision makers). For example, it seems that reflection on a bad decision one has made sometimes greatly decreases the likelihood that one will make similar decisions in the future. Unlike random number generators, many people apparently have the power to learn from their mistakes (which is not to say that random number generators make mistakes). This is one fact about us that can be put to use in developing a response to the problem of present luck. I return to it later.

Perhaps there are two different ways in which being *basically* morally responsible for an action is related to being morally responsible for it. In some cases, it might be that an agent is morally responsible for *A*-ing only if he is basically morally responsible for *A*-ing. There may also be cases in which an agent would be morally responsible for *A*-ing even if he were not basically morally responsible for it, but his being basically morally responsible for *A*-ing contributes positively to the *degree* to which he is morally responsible for *A*-ing. Where might we look for a case of the former kind? An obvious place is the first action for which a given agent is morally responsible. If an agent cannot be morally responsible for any actions he performs unless he is basically morally responsible for at least one action he performs (and if moral responsibility is never retroactive), the first action for which he is morally responsible (if there is one) is one for which he is basically morally responsible, and it is not the case that he would have had some moral responsibility for that action if he had not been basically morally responsible for it.

The question how we can develop from neonates who are not morally responsible for anything into morally responsible agents is an important one in its own right. It certainly merits more philosophical attention than it has received, and reflection on it might improve our chances of finding an attractive answer to the problem of present luck.

In Mele 2006, I developed a response to the problem of present luck that starts on the ground floor, as it were. Perhaps many philosophers who believe that many people are basically morally responsible for some of what they do have not thought much about how ordinary human beings come to be able to act in such a way that they are basically morally responsible for some of their actions. Even so, they realize that ordinary neonates are not morally responsible for anything and that they themselves gradually developed from neonates into relatively sophisticated agents. How does that happen? And where along the way do we begin to perform actions for which we are basically morally responsible (assuming that these philosophers are right in thinking that we sometimes perform such actions)? If you and I are internally indeterministic agents, perhaps that was true of us from the moment we began acting. But people begin acting quite some time before they are morally responsible for any of their actions. How might we develop from tiny indeterministic agents into agents who are basically morally responsible for some of what we do?

Certainly, many parents treat even their four-year-old children as though they regard them as morally responsible for some of the things they do. This obviously does not entail that some four-year-olds do in fact have some moral responsibility for some of their actions. But it suggests that it might be

worthwhile to think about whether ordinary children around that age might be morally responsible for some of their actions. In Mele 2006 (pp. 129–32), I described some relevant shortcomings of normal four-year-olds—compared to normal eight-year-olds (not to mention normal forty-year-olds)—and I intimated that a certain kind of indeterministic agency might not be a greater obstacle to their having some moral responsibility for what they do than these shortcomings are. The shortcomings I highlighted were in impulse control and in capacities for anticipating and understanding the effects of their actions. Part of my aim in my discussion of little agents in Mele 2006 is to move believers in moral responsibility to see the sense in the common idea that moral responsibility comes in degrees and, especially, to see that standards for moral responsibility in young children are plausibly regarded as very modest in comparison with standards for normal adults. If that is plausible, then perhaps it is not outlandish to suggest that some young children are basically morally responsible for some of their actions.

Imagine a universe in which it is known on a planet named Indy that the brain processes that have a direct influence on decisions—including decisions about whether or not to resist temptation—are indeterministic. It is known, as well, that some of these brain processes are associated with shortcomings in impulse control in children who are capable of making decisions. Because these processes are indeterministic, their existence is incompatible with the imagined universe's being deterministic. So their existence is sufficient for the satisfaction of a necessary condition for moral responsibility, if incompatibilism about moral responsibility is true. On Indy, it is also known that indeterministic processes that affect attention have an effect on how well children anticipate the effects of their actions. Fortunately, developmental psychologists on Indy have discovered that good parenting—including encouraging children to behave well, to think about the likely effects of their actions when they are tempted to do something they know is wrong, to take responsibility for their actions, and the like—significantly increases the probability that, over time, children improve markedly in impulse control and in anticipating and understanding the effects of their actions. On Indy, normal children so raised tend to make serious efforts to improve their conduct and they tend to meet with significant success. By the time these children reach adulthood, they have had an important effect on their inclinations, the kinds of decisions they are likely to make as adults, the likelihood that they will resist various temptations, and so on. Even if indeterministic processes in their adult brains still affect their decisions, those processes themselves have been strongly influenced by their past behavior.

The preceding sketch provides a bit of context for what is to come. When pondering whether an indeterministic decision-maker can make a first decision for which he is morally responsible, what sorts of options would it be appropriate to imagine the agent entertaining? If an internally indeterministic adult agent who—because he has never performed a free action (and has never freely omitted to do anything)—has no responsibility at all for the probability at t that shortly thereafter he will decide to rob someone at gunpoint very soon and act accordingly nor for the probability at t that he will instead effectively decide to resist his temptation to commit that crime were to decide on the robbery, the problem of present luck would loom very large.¹⁴ In my discussion of little agents in Mele 2006, I motivate the idea that the problem should be seen as much more modest and tractable when it comes to the relatively trivial first decision for which a normal child might have some moral responsibility. For example, I make the following observation (after quoting from Galen Strawson [2002, p. 451] on heaven-and-hell responsibility):

Obviously, no sane person would think that little Tony [a normal four-year-old] deserves torment in hell—eternal or otherwise—for his bad deeds or heavenly bliss for his good ones. But Tony might occasionally deserve some unpleasant words or some pleasant praise; and, to use Strawson's expression, "it makes sense to propose" that Tony has, for some of his decisions, a degree of moral responsibility that would contribute to the justification of these mild punishments and rewards—even if those decisions are made at times at which the past and the laws leave open alternative courses of action, owing to Tony's being an indeterministic decision maker. (Mele 2006, p. 131)

If, in young children, the bar for basic moral responsibility is pretty low, and low enough that their making an indeterministically caused decision does not preclude their having some degree of moral responsibility for that decision, perhaps an agent's basic moral responsibility can blossom over time into something significantly more robust. In Mele 2006, I developed a libertarian view about how this might happen. One of its planks was a view about how indeterministic agents can learn from their mistakes and successes and shape their probabilities for future action. I wrote:

Probabilities of actions—practical probabilities—for agents are not always imposed on agents. Through their past behavior agents shape present practical probabilities, and in their present behavior they shape

future practical probabilities. The relationship between agents and the probabilities of their actions is very different from the relationship between dice and the probabilities of outcomes of tosses. In the case of dice, of course, the probabilities of future tosses are independent of the outcomes of past tosses. However, the probabilities of agents' future actions are influenced by their present and past actions. (2006, p. 122)

My general strategy was to find a way to make it plausible that a young indeterministic agent might have some very modest basic moral responsibility for some of his actions and then to explain both how moral responsibility can be amplified over time and why, in light of this, the problem of present luck is less threatening to ordinary adult basic moral responsibility than it may initially have seemed to be. If it is assumed that moral responsibility depends on free will (see chapter 5), the connection to free will is made.

Return to the indeterministic agent who was considering robbing someone at gunpoint. In the original story, he has no moral responsibility at all for the pertinent practical probabilities, because he has never acted (nor omitted) freely. In a counterpart story, a lengthy history of actions for which he is morally responsible has done much to shape the probabilities at issue (probabilities with the same numerical values in both cases). Suppose that in both versions of the story, the agent decides on the robbery and executes that decision. Perhaps, dear reader, you are not at all inclined to deny that the agent in the second case is morally responsible for his decision; but if you are, I would bet that your inclination is much weaker in this case than it is in the first case.¹⁵ If my bet is a winner, it seems to matter to you—at least in some cases—how an agent's pertinent practical probabilities were generated. This also would seem to matter to you if you see the second robber, but not the first, as morally responsible for his decision.

Might it be that although the first robber is not morally responsible for his decision, the second one has some basic moral responsibility for his decision that contributes to the degree to which he is morally responsible for it? If this is how things are, then when we think about whether an adult agent is basically morally responsible for a decision he makes, it would be a mistake to ignore how his pertinent practical probabilities at the time were generated. If that is a mistake, seeking a solution to the problem of present luck—as agent causationists do, and as Kane seems to do—solely in nonhistorical properties an agent has at the time of decision (or other action) does not seem promising. It is conceivable that the two robbers are alike in all relevant nonhistorical

respects at the time of decision even though their pertinent practical probabilities were generated in very different ways.

Someone might claim that if young children can be basically morally responsible for indeterministically caused decisions and make such decisions freely, then so can adults, and the problem of present luck was not a problem to begin with. A person who makes this claim may be making the wrong-headed assumption that anything that has a solution was never a problem (see chapter 7, section 7.5). But something else may be going on. The person may be thinking that whatever adequately explains why a normal young child can be basically morally responsible for a decision also adequately explains why the same is true of a normal adult and that the explanation will feature a property that both agents have in common. Of course, part of my own attempted explanation of the possibility of basic moral responsibility in the former case is that the bar for moral responsibility for normal young children is considerably lower than that for normal adults; and, obviously, it is false that the bar for moral responsibility for normal adults is lower than that for normal adults. I agree that if ordinary young children can be basically morally responsible for some decisions they make, then so can normal adults; and I have been suggesting that attention to features of the normal development of normal young agents into normal adult agents helps us see why. Now, if it were to be discovered that whenever we are tempted to act contrary to our better judgment, our practical probabilities of deciding in accordance with that judgment and of succumbing to temptation instead are wholly independent of our past behavior (including decisions and reflection on their consequences), I would concede defeat in this sphere. And if it were to be discovered that our practical probabilities at the time of each and every decision are generated randomly by a mechanism that takes as input only the facts about which options we are entertaining at the time, I would concede general defeat. But there is no reason to believe that such discoveries will be made. In normal adult agents, it may be that moral responsibility for relevant practical probabilities and basic moral responsibility for a particular decision jointly contribute to a level of moral responsibility for that decision that far exceeds what can be found in normal young children.¹⁶

Obviously, it is not age—for example, being four years old rather than forty—that directly matters in my discussion of the behavior, capacities, and practical probabilities of children and adults. Psychological development is the issue, and it is imperfectly correlated with age. There are intellectual and emotional prodigies, adult agents near the other end of the spectrum, and so

on. My discussion of little agents in Mele 2006 was explicitly about normal children in normal circumstances.

I touch on just one further issue before wrapping things up. Recall my observation at the end of section 10.2 that Cathy is not morally responsible for the contrastive fact that she shocked the gray kitten *rather than* the white kitten because she is not morally responsible for the fact that she did not shock the white kitten. Cathy deserves no moral credit for the failure of her effort to shock the white kitten. What about Bob in a story in which what goes on is captured by *DLV*? Might he be morally responsible for deciding to cheat *rather than* deciding to do the right thing (or for the fact that he makes the former decision rather than the latter). Well, if, past intentional, uncom-pelled behavior of his played a significant role in shaping his character and the antecedent probability that he would decide to cheat, and if better behavior was open to him on many relevant occasions in the past, behavior that would have given him a much better chance of deciding to do the right thing on this occasion, then maybe so. But this, as I say, is a long story that I have spun elsewhere. Notice also that, in deciding to cheat, Bob was *deciding against* doing the right thing (see chapter 6, section 6.5). The claim that he freely decided to cheat and freely omitted to decide to do the right thing is in the running, unlike the claim that Cathy freely shocked the gray kitten and freely omitted to shock the white kitten. The claim about Cathy clearly is false; she was trying to shock the white kitten (while also trying to shock the gray kitten).

I mentioned that my error theory is for a limited audience. Some people who regard *DLV* as lacking the resources to provide what is needed for basically free action may require something for free action that no event-causal theory can give them: luck-excluding control over what they choose in a scenario in which it is at no time determined what they will choose.¹⁷ But Kane is not such a person. He means to get by with an event-causal view of action (including choice) production, and he acknowledges the presence of luck in cases of basically free actions (1999b, p. 233, quoted earlier). What I have argued here is that anyone with Kane's aspirations—and any event-causal libertarian—should prefer my daring libertarian view to Kane's concurrent-efforts view.¹⁸

Notes

1. The details of “daring libertarianism” appear in my presentation of what I call “daring soft libertarianism” (see Mele 2006, chap. 5). A *soft* libertarian is open to compatibilism in a certain connection, asserting that “free action and moral

responsibility [may be] compatible with determinism but . . . the falsity of determinism is required for . . . more desirable species of” these things (p. 95). A *daring* libertarian maintains that there are free actions of such a kind that it is at no time determined that the action will occur. A daring soft libertarian endorses both of these theses. Eventually, I make the obvious point that the softness—that is, the openness to compatibilism—can simply be subtracted from daring soft libertarianism (that is, without modifying anything else), yielding what I call “daring libertarianism” (2006, pp. 202–3).

2. The brackets are present in the quoted text. On senses 2 and 3, also see Kane 1996, pp. 77–78.
3. For complications introduced by Frankfurt-style cases and an associated notion of *basically** free action, see Mele 2006, pp. 115–17, 203–5.
4. Also see Kane 1999a, 2000, 2002, and 2011. Readers who balk at the thought that an agent may *try to choose to A* (Kane 1999b, pp. 231, 233–34; 2011, pp. 391–92; 2014, pp. 193–202, 208–9) may prefer to think in terms of an agent’s trying to bring it about that he chooses to *A*.
5. Even if actual people never consciously represent the efforts at issue in this way, Kane can claim that they unconsciously do so.
6. *LD_{Fd}*, in chapter 6, is formulated in terms of directly free decisions. In *LD*, “basically” replaces “directly.”
7. The assertion that Cathy is morally responsible for the fact that she shocked the gray kitten rather than the white kitten—that contrastive fact—should be distinguished from the assertion that Cathy is morally responsible for the fact that she shocked the gray kitten rather than for the fact that she shocked the white kitten.
8. *Directly* free actions are to be distinguished from, for example, free actions of Kane’s type 2 that are deterministically caused by their proximal causes. On a typical libertarian view, all directly free actions are basically free.
9. Here, taking my lead from Kane, I do not treat “voluntary” as entailing “basically free.”
10. On *DLV*, an analogue of a basically free choice is possible in some Frankfurt-style cases. See note 3 for references.
11. On prospective choices as goals, see Kane 2014, pp. 193–94.
12. A novice may suggest that Kane can dramatically improve his view by claiming that one collection of reasons or the other prevails before the choice is made. Imagine a scenario in which Bob’s effort to choose to cheat has the result that at 200 milliseconds (ms) before *t* it is determined that he will choose at *t* to cheat. Imagine also that in another possible world with the same laws of nature and the same past up to *t*–200 ms, Bob’s effort to choose to do the right thing has the result that at *t*–200 ms it is determined that he will choose at *t* to toss the coin straightaway. The problem of present luck has not disappeared; it has been moved back 200 milliseconds.
13. I am grateful to Helen Beebe for recommending that I consider this reply.

14. Here I assume that an agent who has never acted freely and has never freely omitted to do anything is not morally responsible for anything. That assumption may be challenged, of course (see chapter 5).
15. Not all readers should take this bet personally. For example, I would not make this bet with readers who announce their conviction that moral responsibility is possible only in worlds in which determinism is true.
16. Given the definition in play of basic moral responsibility, this obviously depends on the falsity of the following proposition: all of our decisions are deterministically caused by their proximal causes. Assessing that proposition is beyond the scope of this book.
17. I am not suggesting that some other theory can accomplish this trick.
18. Much of this chapter is based on Mele n.d.c, material from which was presented in 2015 at Dartmouth College, the University of Manchester, and the Royal Institute of Philosophy. I am grateful to my audiences for productive discussion. Section 10.5 derives from Mele 2013d, an article on which Randy Clarke and Stephen Kearns provided useful feedback.

Living Without Agent Causation

IN THIS BOOK I have examined several arguments that allegedly falsify event-causal libertarianism or some aspect of my position on it, and I have shown that all of them are unsuccessful. The failed arguments include the *same-control argument* (chapter 7), the *more-control argument* (chapter 7), Derk Pereboom's "disappearing agent objection" (chapter 8), and Randolph Clarke's argument for the falsity of the thesis (*LDFd*) that even if the difference between what an agent decides at *t* in one possible world and what he decides at *t* in another possible world with the same past up to *t* and the same laws of nature is just a matter of luck, the agent may make a directly free decision at *t* in both worlds (chapter 6). These unsuccessful arguments undermine neither *LDFd* nor the standard event-causal libertarian thesis (*LFTe*) that there are indeterministic agents who sometimes act freely when their actions are indeterministically caused by their proximal causes.¹

The failed arguments at issue tend to be especially popular among philosophers who contend, as Clarke and Pereboom do, that having free will depends on having agent-causal powers. In chapter 7, I mentioned Clarke's skepticism about the metaphysical possibility of agent causation and Pereboom's agnosticism about its metaphysical and conceptual possibility. I observed there that if agent causation is impossible, then any bar for free will that includes agent causation is impossibly high, and I made the equally obvious point that the distance between an impossibly high bar and achievable event-causal libertarian and compatibilist bars is incalculable. I suggested that this may help to explain why some philosophers who contend that free will depends on agent causation are unimpressed by differences between event-causal guidance control and event-causal regulative control. As I put the point earlier, because both kinds of control are incalculably distant from any impossible species

of control, they may appear to be so close to each other that any difference between them cannot have an interesting bearing on free will.

In *Free Will and Luck* (Mele 2006, pp. 203–4), I observed that much conceptual reasoning is done under conditions of empirical uncertainty, and I suggested that if some theorists were to come to know certain empirical truths, they might find themselves reasoning differently about conceptual theses they accept; on reflection, they might come around to the view that these theses are false. (If the suggestion seems wildly implausible, I recommend reading Mele 2006, pp. 202–6.) Much conceptual reasoning also happens under conditions of *conceptual* uncertainty, and much metaphysical reasoning takes place under conditions of *metaphysical* uncertainty. Suppose a knockdown proof were to emerge that agent causation is both metaphysically and conceptually impossible, a proof easily understood by educated people. What effect might that have on some philosophers' reasoning about event-causal libertarianism? That is my guiding question in this chapter.

11.1. What Would Be Lost?

With the imagined proof in place, some skeptics about free will who use the proposition that free will depends on agent causation as a plank in their argument for their skepticism would make adjustments. The title of Pereboom 2001, *Living without Free Will*, is the inspiration for this chapter's title. Consider the use Pereboom might make of the imagined proofs. Rather than bother to argue on evidential grounds—controversially, I should add (see Clarke 2010)—that agent causation does not actually exist, Pereboom presumably would cite or rehearse the imagined proof that agent causation is impossible. After all, anything impossible is non-existent, and an easily understood knockdown proof has more clout than a controversial argument. But with agent causation off the table, some other incompatibilists may find the virtues of the positive side of event-causal libertarianism easier to appreciate. When we compare a shiny new Ferrari with a car that is advertised as doing much more, including creating cars with agent-causal powers, the Ferrari may seem like a real dud. If we become convinced that the super car is an impossible car, the Ferrari might start looking very appealing.

What would be lost if agent causation were off the table? Well, what are agent-causal powers supposed to give agents that they cannot otherwise have? A brief review of some issues surrounding some arguments discussed in earlier chapters will prove useful in answering these questions.

1. *The same-control argument.* A premise of this argument is the claim that “the active control that is exercised on [an event-causal libertarian] view is just the same as that exercised on an event-causal compatibilist account” (Clarke 2003, p. 220). Agent-causal powers are supposed to provide for active control that is *different* from what any event-causal compatibilist view can provide. As I pointed out in chapter 7, regulative control is, *by definition*, different from guidance control (or compatibilist control). So the same-control argument fails; it has a false premise. Of course, it may be replied that what is important about agent-causal powers is not only that they are different from compatibilist powers but also that they provide for *more* control than compatibilist powers do and more control than any event-causal libertarian view can dish out. This leads to the next item on my numbered list—the *more-control argument*. But before I turn to it, I repeat a point I made in Mele 2006 and comment briefly on it.

I repeat the point by direct quotation. The paragraph quoted follows on the heels of a critique of a feature of Derk Pereboom’s treatment of agent causation in his 2001 book:

Comparing an agent whose decision is indeterministically caused by events alone with a counterpart agent whose decision is “brought about as characterized by an integrated agent-causal account” ([Clarke] 2003, p. 159), Clarke contends that the latter agent exercised greater active control; he exercised a further power to causally influence which of the open alternatives would come about. In so doing, he was literally an originator of his decision, and neither the decision nor his initiating the decision was causally determined by events. This is why [he] is responsible for his decision, and why it was performed with sufficient active control to have been directly free. If this explanation is correct, then . . . the concept of agent causation is crucially relevant to the problem of free will. (p. 160)

The central point made in the preceding paragraph applies here as well. Are we to believe that the pertinent cross-world difference is not just a matter of luck? If so, why? Are we to believe that although that difference is a matter of luck, the agent is morally responsible for his decision and makes it freely? If so, why? Why should we not believe that the difference is just a matter of luck and that, because it is, the agent is not morally responsible for his decision and does not make it freely, despite exercising “the further power” Clarke cites? I may try to lift a weight using the power of my right arm alone and fail. I may try

again, this time using in addition the further power of my left arm; and I may fail again, the combined powers not being up to the task. If the weight is a ton, the combined powers are not enough to give me even a ghost of a chance of lifting it. For all I have been able to ascertain, the combination of agent causation with indeterministic event causation is similarly inadequate. (Mele 2006, pp. 68–69)

The point made in the quoted paragraph about the integration of agent causation with indeterministic event causation echoes a point I had just made about unintegrated agent causation. It is far from clear how agent causation is supposed to solve the problem of present luck even if it is associated with control that is different from what compatibilists and event-causal libertarians can provide and even if it can be added to event-causal control. The questions I raised about this a decade ago are not rhetorical, and they still have not been answered satisfactorily. And, of course, if agent causation is impossible, it cannot solve the problem.

2. *The more-control argument.* The thrust of this argument is that event-causal libertarianism fails because it does not accommodate any more control than compatibilism does. As I explained in chapter 7, proponents of this argument have neglected to tell us how to measure amounts of control and how to weigh guidance and regulative control on the same scale. Owing to this neglect, the argument is toothless. But what matters for present purposes is that the more-control argument is often paired with the claim that agent-causal powers provide for more control than event-causal powers do. What, exactly, does this extra control enable agents to do? In chapter 8's discussion of Pereboom's disappearing agent objection and some associated work by Meghan Griffith, I mentioned an idea that may have the appearance of being an answer to this question. It is next in line.

3. *Pereboom's disappearing agent objection.* The basic idea here is that in the case of "event-causal libertarian agents" (Pereboom 2014, p. 32) nothing settles whether a decision that is at no time determined to occur will occur and these agents therefore "lack the control required for basic desert moral responsibility" for these decisions (p. 32). Agent-causal powers are supposed to provide for this control. Unfortunately, as I pointed out in chapter 8, Pereboom neglects to provide significant guidance on what he means by an agent's settling whether a decision will occur. I speculated that he might view settling whether one will decide to *A* as something that requires exercising *complete control* over what one does or does not decide, and I explored some potential readings of "complete control." Suppose someone were to propose

the following: *S* exercises complete control over whether a decision to *A* that is at no time determined to occur (in the “deterministic causation” sense of “determined”) will occur only if he exercises his agent-causal power in such a way that at some point before the decision is made there is no chance that his decision to *A* will not occur. Even if satisfying this condition promotes complete control, it also closes the door on something typical libertarians regard as possible in typical cases of directly free decision—namely, that right up to the time at which the agent decided to *A*, there was a chance that he would not decide to *A*. What does the complete control we are looking for look like when this door is left open?

Despite my efforts, I was unable to find in the work of philosophers who regard having agent-causal powers as required for having free will the sort of guidance I would like on what it is for an agent to settle whether a decision will occur or to exercise complete control over whether a decision will occur. In the absence of significant guidance on what this settling and complete control amount to, I find it difficult to worry much that an agent who does not do something called “settling” whether a certain decision will occur and does not exercise something called “complete control” over whether a certain decision will occur does not decide freely. If I were to learn that this settling and exercising are impossible because agent causation is impossible, I would not despair at all. I have never seen myself as having any theoretical use for agent causation. But other philosophers may have a very different reaction to this news; and libertarians who worried that they might need to appeal to agent causation to solve the problem of present luck (or something resembling that problem) but were uneasy about that might see what they regarded as shortcomings of event-causal libertarianism in a new light. If luck-quashing settling whether a certain decision will occur and exercising luck-precluding complete control (which some theorists may regard as the same thing) are impossible in scenarios of the sort that illustrate the problem of present luck, should we believe that, even so, they are required for any decisions made in these scenarios to be directly free? (I have pointed to obscurities in discussions of this settling and complete control. If you, dear reader, are tempted to claim that these things are required for directly free decisions, what do *you* mean by “settling” whether a decision will occur and by exercising “complete control” over whether a decision will occur in scenarios of the sort at issue?) Or does my daring libertarian view, for example, capture all the control that libertarians should believe is needed for directly free decisions? Does the conceptually and metaphysically *possible* control to which this view appeals do all the work that needs to be done on the control

front? And does taking magical control off the table help the virtues of event-causal regulative control shine through?

I highlighted the problem of present luck in the preceding paragraph. Why is that? Recall (from chapter 8) that Pereboom locates his disappearing agent objection in the family of “luck objections” and views it as the member of this family of objections “that reveals the deepest problem for event-causal libertarianism” (2014, p. 32). He maintains that only agent causation is capable of handling the threat that present luck poses for libertarianism—if agent causation is possible, that is. Recall also (from chapter 8) Timothy O’Connor’s claim that “the agent causationist takes it to be a virtue of her theory that it enables her to avoid a ‘problem of luck’ facing other indeterministic accounts” (2011, p. 325). And recall my discussion in chapter 6 of Clarke’s unsuccessful argument against the thesis (*LDFd*) that even if the difference between what an agent decides at *t* in one possible world and what he decides at *t* in another possible world with the same past up to *t* and the same laws of nature is just a matter of luck, the agent may make a directly free decision at *t* in both worlds. Clarke insists that such luck precludes directly free decisions and that agent causation needs to be wheeled in. In his judgment, as I reported, relevant arguments collectively “incline the balance against the possibility of substance causation in general and agent causation in particular” (2003, p. 209). He adds: “We should doubt the possibility of agent causation, but we should not be very certain about the matter” (p. 210). Libertarians who share this cautious, tentative skepticism and libertarians who move beyond it to confidence that agent causation is impossible and therefore cannot solve the problem of present luck have options. One option is my daring libertarian view, including its response to the problem of present luck. Another is trying to develop a superior event-causal libertarian response to that problem. There are other options, of course, including (but not limited to) embracing compatibilism and siding with Pereboom in denying that free will exists.

11.2. Direct Control

Although I have used the word “control” a lot in this chapter, I have not used the expression “direct control” here, an expression that appears occasionally in earlier chapters in my discussions of other people’s work. John Bishop sets the stage for a brief but instructive discussion of direct control by sketching “a

very bad argument" (1989, p. 70) that he attributes to Peter van Inwagen. The argument runs as follows:

Consider a mechanism that on the push of a button flashes either a red light (with 50 percent probability) or a green (with 50 percent probability). Then an agent who pushes the button has no control over which light flashes. Now if the mental states that constitute an agent's reasons caused matching behavior only probabilistically, by parity of reasoning, the agent would have no control over whether that matching behavior was produced or not. Hence, actions cannot be behavior with merely probabilistic mental causes. (Bishop 1989, p. 70)

Bishop objects that the argument draws "a false analogy" between a light's flashing and an action (p. 70). Something true of a probabilistically caused nonaction might not be true of a probabilistically caused action. And although the agent only "indirectly controls" whether the light flashes or not, his pressing the button—an action—is "an exercise of direct control" (p. 71). In this way, Bishop contends, at least some actions (for example, John's pressing a button) differ from outcomes of an action (for example, a light's flashing). In Bishop's view, whenever agents act, they exercise control (pp. 23, 25)—and, more specifically, direct control. Clarke agrees: "In every instance of action, the agent exercises some degree of direct active control" (2003, p. 76). Bishop's event-causal theory of action is meant to accommodate this idea.

As Bishop understands direct control, it is nothing out of the ordinary. That is reassuring. But I, at least, am left with some questions. Consider the following two assertions: Joe exercised direct control; Joe's pressing the button was an exercise of direct control. They have the ring of incompleteness. "Over what?" one wants to ask. If Joe raised his right arm (in an ordinary way) to vote for a motion at a meeting, we might say that he exercised direct control over his arm—or over how his arm moved, or over the motions of his arm. If he tied a rope around his right arm and then raised his right arm by raising the rope with his left arm, did he exercise direct control over his left arm and indirect control over his right arm? Did he exercise direct or indirect control over the rope? If Joe tied a rope around a log and then pulled the log to a woodpile, did he exercise direct control over the log? Did he exercise direct control over the rope and indirect control over the log? Did whatever he exercised direct control over extend no further than his body?

All these questions and many more would be answered by a full account of direct control. To the best of my knowledge, nothing approaching a full account of it exists, and I do not try to provide one here. Although the project of developing a full account would be interesting, I believe that many of the details would not have a special bearing on free will. Even so, I do not know how to proceed in a discussion of direct control without offering readers more guidance on what it might be than I have encountered in my own reading.

Suppose, following Bishop and Clarke, that whenever agents act intentionally, they exercise direct control over something or other. In an ordinary case of raising one's right hand over one's head, what does the agent exercise direct control over? One seemingly reasonable answer is *motions of his arm and hand*. What does an agent exercise direct control over when he decides to *A*? An analogous answer, given the account of deciding developed in chapter 2, is *his acquisition of an intention to A*.

Consider an alternative answer to my question about an ordinary case of raising one's right hand. Someone may claim that the agent exercises direct control over his raising his right hand—that action. The answer in the preceding paragraph represents the action at issue—the hand-raising—as an exercise of direct control over something. And the answer in the present paragraph represents that very action as something the agent exercises direct control over. According to one view of things, we *perform* our actions and whenever we act intentionally we exercise some direct control over one or more nonactions—for example, motions of our bodies or our acquisition of an intention. But according to the claim at issue now, whenever we act intentionally we exercise some direct control over some action.² This latter idea, Bishop claims, is wrongly attributed to agent causationists. He asserts that a certain alleged difficulty for agent causationists dissolves once a misunderstanding of their view is exposed: “The theory is that actions consist in the causing by their agents *of certain events or states of affairs*. Thus, agents are not held to agent-cause their *actions* . . . but rather the events or states of affairs that are, so to say, *intrinsic* to their actions” (1989, p. 68). If, according to agent causationists, agents exercise direct control only over what they agent-cause, then, if Bishop's interpretation is correct, they do not exercise direct control over their actions.

I paired the idea that when Joe raises his hand in the normal way he exercises direct control over motions of his arm and hand with the idea that when Joe decides to *A* he exercises direct control over his acquisition of an intention to *A*. Having a label for the conception of direct control reflected in these ideas will facilitate discussion. I will say that, when combined with the idea that direct control is never exercised over actions, they reflect a *Bishop-style* conception. The alternative idea that when Joe raises his hand he exercises direct control

over that action is paired with the idea that when Joe decides to *A* he exercises direct control over his deciding to *A*—that action. These alternative ideas reflect what I dub a *Knight-style* conception of direct control (a name chess enthusiasts might like; *Gangnam-style* was an alternative I briefly considered).³

If I were to offer an account of something I call “exercising direct control,” I would want the term “direct” to do important work. The following certainly seems to be a plausible requirement on an agent’s exercising *direct* control over *X*: If *S* exercises direct control over *X*, then *S* does not exercise control over *X* only by exercising control over something else (or, more precisely, something that does not include *X*).⁴ If this proposition about direct control is true, then if Joe exercises control over the log only by exercising control over the rope, he does not exercise direct control over the log. And if he exercises control over the rope only by exercising control over relevant bodily motions, then he does not exercise direct control over the rope. Comparable examples framed in a Knight-style way would involve an agent’s exercising control over one action only by exercising control over another action (that does not include the former action). If Joe exercises control over his moving of the log only by exercising control over his pulling of the rope, he does not exercise direct control over the former action.⁵ If asked, some proponents of a Knight-style conception of direct control might assert that agents have direct control only over their *basic* actions.⁶ Similarly, some proponents of a Bishop-style conception might assert that agents have direct control only over nonactions “that are, so to say, *intrinsic* to” their basic actions (Bishop 1989, p. 68).⁷ (Example: The rising of one’s arm is intrinsic to one’s raising it.)

I have supposed, following Bishop and Clarke, that whenever agents act intentionally, they exercise direct control over something or other. (To avoid excessive repetition, I dub this supposition *S*.) If there are basic actions, an agent performs at least one basic action whenever he acts intentionally. In light of this, a pair of suggestions I just made may be augmented. A proponent of a Knight-style conception of direct control might assert that agents exercise direct control over *all and only* their basic actions and a proponent of a Bishop-style conception might assert that agents exercise direct control over *all and only* nonactions that are intrinsic to their basic actions. Readers should treat these “all and only” ideas as working assumptions.

An obvious point should be made. Unless all basic actions are free actions, the Bishop-style and Knight-style views of direct control, as thus far developed, leave a considerable gap between direct control and free will. Of course, this is what one should have expected, given supposition *S*. With that supposition in place, unless it is true that whenever agents act intentionally they do something or other freely, there will be times at which agents exercise direct

control and do not act freely. And there are many scenarios in which, although an agent does something intentionally, he does not do anything freely. For just one kind of example among many, consider very young children who have developed the capacity for intentional action but have not developed the capacity for free action. Also, if there are possible worlds in which lots of agents frequently act intentionally but no agent has free will, then, given supposition *S*, there are possible worlds in which agents often exercise direct control but no agent ever acts freely. I leave it to readers to identify such worlds for themselves. An incompatibilist can pick any deterministic world in which there are beings that often act intentionally, unless he or she believes that intentional action is incompatible with determinism. People who hold that there are nonhuman animals that act intentionally but are incapable of acting freely can pick a world in which those animals are the only agents. Finding a world you regard as illustrating the point is left up to you, dear reader.⁸ However the directness of direct control is to be explicated, some readers may worry that Bishop-style direct control leaves important facts about our control out in the cold. Don't we *have* some control over whether we decide to *A* or decide to *B*, in some cases? And when we decide to *B* in such cases, don't we *exercise* direct control over our deciding to *B*? These are among the questions such readers may ask.

I offer some potential answers from a Bishop-style perspective. The control we have over whether we decide to *A* or decide to *B* in typical cases of the sort at issue consists partly in our being able to decide to *A* for reasons that recommend *A*-ing and able to decide to *B* for reasons that recommend *B*-ing. And there is plenty of room for exercises of indirect control over whether we decide to *A* or decide to *B*. Consider a representative case of decision-making that is informed by evidence gathering and thoughtful reflection. Joe will be moving to another state soon to take a new job, and he is thinking about whether to buy a house there soon or rent for a while and buy a house later. He gathers information about houses and real estate agents, asks a real estate agent to show him around, looks at houses, gathers more information, and so on. These actions have an effect on what he eventually decides to do—and does—about his housing situation. In performing them, he exercises indirect control over what he will decide. In the end, he decides to buy house *H*. In performing actions of the kind mentioned, he exercises indirect control over his deciding to do that, over his doing it, over whether he decides to buy house *H* or decides to do something else instead, and over whether he buys that house or instead does something else.

I commented on exercising control that we have over whether we decide to *A* or decide to *B*. But my comments were about indirect control. What

about direct control? On a Bishop-style view, when we decide to *B*, we do exercise direct control; we exercise it over our acquisition of an intention to *B*. And if we had decided to *A*, we also would have exercised direct control—in this case, over our acquisition of an intention to *A*. In typical cases in which an agent with the pair of abilities I mentioned continues to regard both *A* and *B* as live options at the moment of decision, a Bishop-style theorist may claim that in exercising direct control over his acquisition of the intention he acquires, he exercises direct control over *whether*, at the time, he acquires an intention to *A* or acquires an intention to *B*.⁹

I suggested that a proponent of a Bishop-style conception of direct control might assert that agents exercise direct control over all and only nonactions that are intrinsic to their basic actions. The suggestion now is that some such exercises amount to exercises of direct control by an agent over whether, at the time, he acquires an intention to *A* or acquires an intention to *B*. If it is assumed that the abilities at issue are what in chapter 4 I called *O*-abilities, the discussion in section 6.5 of chapter 6 is directly relevant here. (Recall that an agent who did not do *A* at *t* was *O*-able at the pertinent time to *A* at *t* only if in a possible world with the same laws of nature and the same past up to *t*, he *A*-s at *t*.)

In that section of chapter 6, I explained that, for all anyone has shown, it may have been up to an agent who decided at *t* to *A* whether he would decide then to *A* or decide then to *B* even if the difference at *t* between a world in which he decides at *t* to *A* and a world with the same past up to *t* and the same laws in which he decides at *t* to *B* is just a matter of luck. Similarly, for all anyone has shown, this agent may have exercised direct control over whether he acquired an intention to *A* or acquired an intention to *B*, even if the difference just mentioned is just a matter of luck. Perhaps something that some theorists may call “complete control” is such that it is exercised in a scenario of the sort at issue only if an impossible power is at work. But exercising direct control over *X* does not entail exercising complete control over *X* on any Bishop-style or Knight-style conception of direct control according to which, whenever agents act intentionally, they exercise direct control over something.

To illustrate my claim about the relationship between direct control and complete control, I draw on a pair of examples from chapter 8. Recall the scenario in which Sol is using a keyboard with a randomizer on it that ensures that there is always a small chance that a key he is trying to press will stick and fail to make contact with the switch under it. Just now, Sol fully depressed the Q key, which is what he was trying to do. On some accounts of basic action, his doing so is a basic action. And proponents of Bishop-style and Knight-style views who regard it as such (as I am interpreting such views) will say

that Sol exercised direct control over the key's contacting the switch (Bishop-style) or over his fully depressing the key (Knight-style). That Sol obviously does not exercise complete control over these things is no problem at all for their claims.

A theorist who is more restrictive about basic actions may say that Sol's basic action in this case is his moving a finger in a certain way or his trying to depress the key fully. Theorists who take this position will say that Sol exercised direct control over certain bodily motions (Bishop-style) or over the finger-moving action or the attempt (Knight-style). In a second example from chapter 8 (slightly modified), a randomizer has been installed in Sol's brain that ensures that there is always a small chance that his proximal intentions to press a key will not be followed by a corresponding attempt nor by any relevant bodily motions. Just now, Sol fully depressed the key; and, of course, he lacked complete control over whether he would try to press the key on this occasion and over whether he would move his finger. But that fact is entirely consistent with the truth of the claims at issue about what he exercised direct control over.

Proponents of a Bishop-style conception of direct control have questions of their own to raise. According to a Knight-style conception, what is it to exercise direct control over one's deciding to *A*? What is there to this exercise beyond the agent's deciding to *A* or his deciding to *A* for relevant reasons? And, in a case in which an agent continues to regard both *A* and *B* as live options at the moment of decision and is *O*-able at the time to decide to *A* then and *O*-able at the time to decide to *B* then, what is it for him to exercise direct control over whether he decides then to *A* or decides then to *B* beyond deciding one way or the other (for relevant reasons) at that time?

A Knight-style theorist will not claim that an agent's exercise of direct control over his deciding to *A* consists in some distinct action over which he exercises some control. For in that case, the control the agent exercises over his deciding to *A* would be indirect: he would be exercising control over his deciding to *A* only by exercising control over some distinct action. Perhaps, then, our Knight-style theorist will claim that the agent's exercising direct control over his deciding to *A* just is his deciding to *A*. If the answer is supposed to apply to all possible instances of deciding to act, it can be stated as follows: Necessarily, for any *A*, an agent's deciding to *A* is identical with his exercising direct control over his deciding to *A*. (Notice that if exercising direct control over one's deciding to *A* is understood in this way, it suffices for freely deciding to *A* only if all possible decisions are free.) Someone who finds

this answer attractive may take a parallel position on basic action: necessarily, any basic action, *A*, is identical with the agent's exercising direct control over *A*.

One can imagine a proponent of a Knight-style view who is a fan of agent causation and contends that to exercise direct control over one's deciding to *A* is precisely to agent-cause one's deciding to *A* and to exercise direct control over whether one decides to *A* or decides to *B* in a situation of the sort described in the preceding paragraph is precisely to agent-cause one action of decision-making or the other. This claim would raise questions that have what is by now a very familiar ring. What is it about agent causation in virtue of which agent-causing one's deciding to *A*, thereby exercising direct control over one's deciding to *A* and over whether one decides to *A* or decides to *B*, provides a solution to the problem of present luck? Why is the featured difference at *t* between a possible world in which Bob agent-causes at *t* a decision to cheat, thereby exercising direct control over his deciding to cheat, and another possible world with the same past up to *t* and the same laws of nature in which Bob agent-causes at *t* a decision not to cheat, thereby exercising direct control over his deciding not to cheat, not a matter of luck? Why is the difference in how Bob employs his power of direct control at the time not a matter of luck? One may stipulate that agent causal power enables one to decide in a luck-precluding way in scenarios of the sort at issue, but *stipulating* this is uninteresting. And, to return to the theme of this chapter, if agent causation is impossible, libertarians may take comfort in the idea that we have decision-making *O*-abilities of the kind I mentioned and Bishop-style direct control over intentions that we acquire.¹⁰

I have been speculating about how some libertarians might proceed in the face of a newfound conviction that agent causation is impossible. Why, a reader might ask, don't I argue for the impossibility of agent causation? That is a fair question. Part of the reason is that there are different conceptions or accounts of agent causation on offer, and I am disinclined to sort through all of them and defend a judgment about each one. I stick with the position I defended on agent causation in Mele 2006: namely, that it does not provide a solution to the problem of present luck. Libertarians who accept that position have little incentive to appeal to agent causation. The good news for libertarians, if I am right, is that the problem of present luck can be adequately dealt with from an event-causal libertarian perspective.

Why, then, am I not a libertarian? As I have reported, I am not persuaded by any argument for incompatibilism about free will; and such incompatibilism

is a defining feature of libertarianism.¹¹ I continue to keep compatibilism on the table. I have not said much about compatibilism here, but there will be other occasions for that. As I have also reported (chapter 10), I do not know of strong evidence for the proposition that human brains have all the indeterministic features they would need to have for a theoretically attractive event-causal libertarian view to be true. But I do not know of strong evidence for the falsity of this proposition either. The positive side of a view like my daring libertarianism is in the running. Science has not closed the door on it.¹² Nor, of course, has that door been closed by the unsuccessful arguments against event-causal libertarianism examined in this book.¹³

Notes

1. *LFTe* is an explicitly event-causal version of *LFT* (see chapters. 7 and 9).
2. Someone who makes this claim may also claim that we exercise direct control over some bodily motions and other nonactions. And, for that matter, for all that has been said so far, someone who offers the answers mentioned in the preceding paragraph may also claim that we exercise direct control over some of our actions.
3. When alternatives to these two conceptions of direct control are articulated, they can be assessed. My aim here does not include identifying all possible conceptions of direct control. One might consider developing a conception of direct control over overt actions that assumes that such actions are to be identified with bodily motions that are caused in certain ways and, in the case of many actions, have certain effects. (And one might use this conception as a model for dealing with purely mental actions.) For a powerful critique of this way of viewing overt actions, see Hornsby 1980, chap. 1.
4. Depending on how one understands “by,” one may confidently assert that *S* may exercise direct control over *X* by exercising direct control over *X* and *Y*; and what is referred to by “*X* and *Y*” is something other than *X*. Hence the parenthetical clause.
5. Readers should bear in mind my policy on action individuation in this book (see chapter 1, section 1.1). To forestall confusion, I observe that a coarse-grained theorist about action individuation would put the point this way: If Joe exercises control over his action under the description “moves the log” only by exercising control over his action under the description “pulls the rope,” he does not exercise direct control over his action under the former description.
6. A coarse-grained theorist about action individuation who has a use for a notion of basic action would say that the same action may be basic under some descriptions and nonbasic under other descriptions.
7. Bishop himself, in a discussion of a scenario in which an agent presses a button to make a light flash, asserts that “the agent would directly control the button” (1989, p. 71). He may treat pressing a button as a basic action in many ordinary cases. In

Bishop's view, "basic actions are those the agent can perform directly, without having to find other means for their achievement" (p. 128). In a discussion of a well-known story spun by Donald Davidson, Bishop refers to a certain agent's "letting go of the rope" as a basic action (p. 133).

8. An anonymous referee of Mele n.d.a suggested, reasonably, that some incompatibilists might not accept *S* and that some "authors may unwittingly use the notion of 'direct control' more or less stringently in different dialectical contexts." Philosophers who reject *S* either have or are seeking a conception of direct control that is more exclusive than anything that can encompass *S*. They have the option of understanding my sketches as having something they might label *direct control lite* as their subject matter. They may conceive of those sketches as needing to be augmented with something for which the rejection of *S* may clear the way—something that would secure an action's being directly *free*. The mode of augmentation I have in mind would build on a connection I have explored between exercising direct control and performing a basic action. I have no objection to this way of proceeding. In fact, given that they reject *S*, I recommend it. (Philosophers who accept *S* may regard what these other philosophers are after as something that should be labeled *direct control plus*.) Philosophers who unwittingly use "direct control" to mean different things at different times may benefit from learning that they have done this, but I have no desire to call anyone out on this.
9. In atypical cases, one of the abilities may be lost before the decision is made.
10. The characterizations I presented of exercising direct control over whether (*DC1*) one acquires an intention to *A* or an intention to *B* (Bishop-style) and over whether (*DC2*) one decides to *A* or decides to *B* (Knight-style) involve the agent's exercising direct control over something specific (e.g., his acquisition of an intention to *A* or his deciding to *A*). If and when alternative characterizations of exercising direct control over whether *DC1* and over whether *DC2* are articulated, they can be assessed.
11. Among the views about free will that I have floated are some "soft libertarian" views (see chapter 10, n. 1). They are not genuinely libertarian, owing to their openness to compatibilism.
12. For discussion of some relevant empirical matters, see Mele 2009.
13. Section 11.2 of this chapter derives from Mele n.d.a, I am grateful to Randy Clarke for written comments on that article and to Stephen Kearns for discussion.

References

- Adams, Frederick. 1986. "Intention and Intentional Action: The Simple View." *Mind and Language* 1: 281–301.
- Adams, Frederick, and A. Mele. 1992. "The Intention/Volition Debate." *Canadian Journal of Philosophy* 22: 323–38.
- Anscombe, G. E. M. 1963. *Intention*. 2nd ed. Ithaca, N.Y.: Cornell University Press.
- Armstrong, David. 1980. *The Nature of Mind*. Ithaca, N.Y.: Cornell University Press.
- Audi, Robert. 1979. "Weakness of Will and Practical Judgment." *Notus* 13: 173–96.
- . 1993. *Action, Intention, and Reason*. Ithaca, N.Y.: Cornell University Press.
- Austin, John. 1970. "Ifs and Cans." In J. Urmson and G. Warnock, eds. *Philosophical Papers*, 205–32. Oxford: Oxford University Press.
- Ayer, Alfred. 1954. "Freedom and Necessity." In A. Ayer, *Philosophical Essays*, 271–84. London: Macmillan.
- Balaguer, Mark. 2014. "Replies to McKenna, Pereboom, and Kane." *Philosophical Studies* 169: 71–92.
- Baron, Robert. 1997. "The Sweet Smell of . . . Helping: Effects of Pleasant Ambient Fragrance on Prosocial Behavior in Shopping Malls." *Personality and Social Psychology Bulletin* 23: 498–503.
- Berofsky, Bernard. 2012. *Nature's Challenge to Free Will*. Oxford: Oxford University Press.
- Bishop, John. 1989. *Natural Agency*. Cambridge, Mass.: Cambridge University Press.
- Brand, Myles. 1984. *Intending and Acting*. Cambridge, Mass.: MIT Press.
- Bratman, Michael. 1984. "Two Faces of Intention." *Philosophical Review* 93: 375–405.
- . 1987. *Intention, Plans, and Practical Reason*. Cambridge, Mass.: Harvard University Press.
- Brembs, Bjorn. 2011. "Towards a Scientific Concept of Free Will as a Biological Trait: Spontaneous Actions and Decision-Making in Invertebrates." *Proceedings of the Royal Society Biological Sciences* 278: 930–39.
- Campbell, Charles. 1957. *On Selfhood and Godhood*. London: Allen and Unwin.

- Cappelen, Herman. 2012. *Philosophy without Intuitions*. Oxford: Oxford University Press.
- Child, William. 1994. *Causality, Interpretation, and the Mind*. Oxford: Oxford University Press.
- Chisholm, Roderick. 1966. "Freedom and Action." In K. Lehrer, ed. *Freedom and Determinism*, 11–44. New York: Random House.
- Clarke, Randolph. 2000. "Modest Libertarianism." *Philosophical Perspectives* 14: 21–46.
- . 2003. *Libertarian Accounts of Free Will*. Oxford: Oxford University Press.
- . 2004. "Reflections on an Argument from Luck." *Philosophical Topics* 32: 47–64.
- . 2005. "Agent Causation and the Problem of Luck." *Pacific Philosophical Quarterly* 86: 408–21.
- . 2010. "Are We Free to Obey the Laws?" *American Philosophical Quarterly* 47: 389–401.
- . 2011. "Alternatives for Libertarians." In R. Kane, ed. *The Oxford Handbook of Free Will*, 2nd ed., 329–48. New York: Oxford University Press.
- Dancy, Jonathan. 2000. *Practical Reality*. Oxford: Oxford University Press.
- Davidson, Donald. 1980. *Essays on Actions and Events*. Oxford: Clarendon Press.
- . 1985. "Replies to Essays I–IX." In B. Vermazen and M. Hintikka, eds. *Essays on Davidson*, 195–229. Oxford: Clarendon Press.
- Dennett, Daniel. 1978. *Brainstorms*. Montgometry, Vt.: Bradford Books.
- D'Oro, Giuseppina. 2007. "Two Dogmas of Contemporary Philosophy of Action." *Journal of the Philosophy of History* 1: 10–24.
- Fischer, John. 1994. *The Metaphysics of Free Will*. Cambridge, Mass.: Blackwell.
- . 2006. *My Way: Essays on Moral Responsibility*. Oxford: Oxford University Press.
- . 2012. *Deep Control: Essays on Free Will and Value*. New York: Oxford University Press.
- . 2014. "Toward a Solution to the Luck Problem." In D. Palmer, ed. *Libertarian Free Will*, 52–68. Oxford: Oxford University Press.
- Fischer, John, and M. Ravizza. 1992. "When the Will Is Free." *Philosophical Perspectives* 6: 423–51.
- . 1998. *Responsibility and Control: A Theory of Moral Responsibility*. New York: Cambridge University Press.
- Frankfurt, Harry. 1969. "Alternate Possibilities and Moral Responsibility." *Journal of Philosophy* 66: 829–39.
- . 1988. *The Importance of What We Care About*. Cambridge: Cambridge University Press.
- Gazzaniga, Michael. 2011. *Who's in Charge? Free Will and the Science of the Brain*. New York: Ecco.
- Ginet, Carl. 1990. *On Action*. Cambridge: Cambridge University Press.
- . 2008. "In Defense of a Non-Causal Account of Reasons Explanations." *Journal of Ethics* 12: 229–37.

- . 2014. "Can an Indeterministic Cause Leave a Choice up to the Agent?" In D. Palmer, ed. *Libertarian Free Will*, 37–51. Oxford: Oxford University Press.
- Goldman, Alvin. 1970. *A Theory of Human Action*. Englewood Cliffs, N.J.: Prentice-Hall.
- Greene, Joshua, and J. Cohen. 2004. "For the Law Neuroscience Changes Nothing and Everything." *Philosophical Transactions of the Royal Society B: Biological Sciences* 359: 1775–85.
- Griffith, Meghan. 2010. "Why Agent-Caused Actions Are Not Lucky." *American Philosophical Quarterly* 47: 43–56.
- Grünbaum, Adolph. 1971. "Free Will and the Laws of Human Behavior." *American Philosophical Quarterly* 8: 299–317.
- Haji, Ishtiyaque. 2004. "Active Control, Agent-Causation and Free Action." *Philosophical Explorations* 7: 131–48.
- Hornsby, Jennifer. 1980. *Actions*. Boston: Routledge and Kegan Paul.
- . 1993. "Agency and Causal Explanation." In J. Heil and A. Mele, eds. *Mental Causation*, 161–88. Oxford: Oxford University Press.
- Hume, David. 1739. *A Treatise of Human Nature*. Reprinted in Lewis Selby-Bigge, ed. *A Treatise of Human Nature*. Oxford: Clarendon Press, 1975.
- Jahanshahi, Marjan, I. Jenkins, R. Brown, C. Marsden, R. Passingham, and D. Brooks. 1995. "Self-Initiated Versus Externally Triggered Movements." *Brain* 118: 913–33.
- James, William. [1890] 1981. *The Principles of Psychology*. Vol. 2. Cambridge, Mass.: Harvard University Press.
- Kane, Robert. 1989. "Two Kinds of Incompatibilism." *Philosophy and Phenomenological Research* 50: 219–54.
- . 1996. *The Significance of Free Will*. New York: Oxford University Press.
- . 1999a. "On Free Will, Responsibility and Indeterminism: Responses to Clarke, Haji, and Mele." *Philosophical Explorations* 2: 105–21.
- . 1999b. "Responsibility, Luck and Chance: Reflections on Free Will and Indeterminism." *Journal of Philosophy* 96: 217–40.
- . 2000. "Responses to Bernard Berofsky, John Martin Fischer, and Galen Strawson." *Philosophy and Phenomenological Research* 60: 157–67.
- . 2002. "Some Neglected Pathways in the Free Will Labyrinth." In R. Kane, ed. *The Oxford Handbook of Free Will*, 406–37. New York: Oxford University Press.
- . 2008. "Three Freedoms, Free Will, and Self-Formation: A Reply to Levy and Other Critics." In N. Trakakis and D. Cohen, eds. *Essays on Free Will and Moral Responsibility*, 142–62. Newcastle-upon-Tyne: Cambridge Scholars Publishing.
- . 2011. "Rethinking Free Will: New Perspectives on an Ancient Problem." In R. Kane, ed. *The Oxford Handbook of Free Will*. 2nd ed., 381–404. New York: Oxford University Press.
- . 2014. "New Arguments in Debates on Libertarian Free Will: Responses to Contributors." In D. Palmer, ed. *Libertarian Free Will*, 179–214. Oxford: Oxford University Press.

- Kapitan, Tomis. 1996. "Incompatibilism and Ambiguity in the Practical Modalities." *Analysis* 56: 102–10.
- Kaufman, Arnold. 1966. "Practical Decision." *Mind* 75: 25–44.
- Kavka, Gregory. 1983. "The Toxin Puzzle." *Analysis* 43: 33–36.
- Kearns, Stephen, and A. Mele. 2014. "Have Compatibilists Solved the Luck Problem for Libertarians?" *Philosophical Inquiries* 2: 9–36.
- Lau, Hakwan, R. Rogers, N. Ramnani, and R. Passingham. 2004. "Willed Action and Attention to the Selection of Action." *Neuroimage* 21: 1407–14.
- Levy, Neil. 2011. *Hard Luck: How Luck Undermines Free Will and Moral Responsibility*. Oxford: Oxford University Press.
- Lewis, David. 1986. *Philosophical Papers*, vol. 2. New York: Oxford University Press.
- Markosian, Ned. 1999. "A Compatibilist Version of the Theory of Agent Causation." *Pacific Philosophical Quarterly* 80: 257–77.
- Mavrodes, George. 1986. "Religion and the Queerness of Morality." In R. Audi and W. Wainwright, eds. *Rationality, Religious Belief, and Moral Commitment*, 213–26. Ithaca, N.Y.: Cornell University Press.
- Maye, Alexander, C. Hsieh, G. Sugihara, and B. Brembs. 2007. "Order in Spontaneous Behavior." *PLoS ONE* 2, e443: 1–14.
- McCann, Hugh. 1975. "Trying, Paralysis, and Volition." *Review of Metaphysics* 28: 423–42.
- . 1986a. "Intrinsic Intentionality." *Theory and Decision* 20: 247–73.
- . 1986b. "Rationality and the Range of Intention." *Midwest Studies in Philosophy* 10: 191–211.
- . 1991. "Settled Objectives and Rational Constraints." *American Philosophical Quarterly* 28: 25–36.
- . 1998. *The Works of Agency*. Ithaca, N.Y.: Cornell University Press.
- McGinn, Colin. 1982. *The Character of Mind*. Oxford: Oxford University Press.
- McKenna, Michael. 2001. "Source Incompatibilism, Ultimacy, and the Transfer of Non-responsibility." *American Philosophical Quarterly* 38: 37–52.
- Mele, Alfred. 1987. *Irrationality: An Essay on Akrasia, Self-Deception, and Self-Control*. New York: Oxford University Press.
- . 1992a. *Springs of Action: Understanding Intentional Behavior*. New York: Oxford University Press.
- . 1992b. "Intentions, Reasons, and Beliefs: Morals of the Toxin Puzzle." *Philosophical Studies* 68: 171–94.
- . 1995a. *Autonomous Agents: From Self-Control to Autonomy*. New York: Oxford University Press.
- . 1995b. "Effective Deliberation about What to Intend: Or Striking It Rich in a Toxin-Free Environment." *Philosophical Studies* 79: 85–93.
- . 1995c. "Motivation: Essentially Motivation-Constituting Attitudes." *Philosophical Review* 104: 387–423.
- . 1996. "Soft Libertarianism and Frankfurt-Style Scenarios." *Philosophical Topics* 24: 123–41.

- . 1997. "Agency and Mental Action." *Philosophical Perspectives* 11: 231–49.
- . 2000. "Deciding to Act." *Philosophical Studies* 100: 81–108.
- . 2003a. *Motivation and Agency*. New York: Oxford University Press.
- . 2003b. "Agents' Abilities." *Noûs* 37: 447–70.
- . 2005. "Libertarianism, Luck, and Control." *Pacific Philosophical Quarterly* 86: 395–421.
- . 2006. *Free Will and Luck*. New York: Oxford University Press.
- . 2007a. "Free Will and Luck: Reply to Critics." *Philosophical Explorations* 10: 195–210.
- . 2007b. "Persisting Intentions." *Noûs* 41: 735–57.
- . 2007c. "Reasonology and False Beliefs." *Philosophical Papers* 36: 91–118.
- . 2009. *Effective Intentions*. New York: Oxford University Press.
- . 2010. "Teleological Explanations of Actions: Anticausalism vs. Causalism." In J. Aguilar and A. Buckareff, eds. *Causing Human Actions: New Perspectives on the Causal Theory of Action*, 183–98. Cambridge, Mass.: MIT Press.
- . 2012a. *Backsliding: Understanding Weakness of Will*. New York: Oxford University Press.
- . 2012b. "Intentional, Unintentional, or Neither? Middle Ground in Theory and Practice." *American Philosophical Quarterly* 49: 369–79.
- . 2013a. "Actions, Explanations, and Causes." In G. D'Oro and C. Sandis, eds. *Reasons and Causes: Causalism and Anti-Causalism in the Philosophy of Action*, 160–74. Basingstoke: Palgrave Macmillan.
- . 2013b. "Is What You Decide Ever up to You?" In I. Haji and J. Caouette, eds. *Free Will and Moral Responsibility*. Newcastle-upon-Tyne: Cambridge Scholars Publishing: 74–97.
- . 2013c. "Libertarianism and Human Agency." *Philosophy and Phenomenological Research* 87: 72–92.
- . 2013d. "Moral Responsibility and the Continuation Problem." *Philosophical Studies* 162: 237–55.
- . 2014a. "Kane, Luck, and Control: Trying to Get by without Too Much Effort." In D. Palmer, ed. *Libertarian Free Will*, 37–51. Oxford: Oxford University Press.
- . 2014b. "Luck and Free Will." *Metaphilosophy* 45: 543–57.
- . 2015a. "Free Will and Moral Responsibility: Does Either Require the Other?" *Philosophical Explorations* 8: 297–309.
- . 2015b. "Libertarianism, Compatibilism, and Luck." *Journal of Ethics* 19: 1–21.
- . 2015c. "Luck, Control, and Free Will: Answering Berofsky." *Journal of Philosophy* 112: 337–55.
- . n.d.a. "Direct Control." *Philosophical Studies* (forthcoming).
- . n.d.b. "On Pereboom's Disappearing Agent Argument." *Criminal Law and Philosophy* (forthcoming).
- . n.d.c. "Two Libertarian Theories: Or Why Event-Causal Libertarians Should Prefer My Daring Libertarian View to Robert Kane's View." *Philosophy* (forthcoming).

- Mele, Alfred, and P. Moser. 1994. "Intentional Action." *Noûs* 28: 39–68.
- Mele, Alfred, and D. Robb. 1998. "Rescuing Frankfurt-Style Cases." *Philosophical Review* 107: 97–112.
- . 2003. "BBs, Magnets and Seesaws: The Metaphysics of Frankfurt-Style Cases." In D. Widerker and M. McKenna, eds. *Moral Responsibility and Alternative Possibilities: Essays on the Importance of Alternative Possibilities*, 127–38. Burlington, Vt.: Ashgate.
- Mellor, D. H. 1995. *The Facts of Causation*. London: Routledge.
- Mickelson, Kristin. 2015. "The Zygote Argument Is Invalid: Now What?" *Philosophical Studies* 172: 2911–29.
- Mill, John Stuart. 1979. *An Examination of Sir William Hamilton's Philosophy*. John Robson, ed. Toronto: Routledge and Kegan Paul.
- Miller, Jeff, and W. Schwarz. 2014. "Brain Signals Do Not Demonstrate Unconscious Decision Making: An Interpretation Based on Graded Conscious Awareness." *Consciousness and Cognition* 24: 12–21.
- Montague, P. Read. 2008. "Free Will." *Current Biology* 18: R584–85.
- Nagel, Thomas. 1986. *The View from Nowhere*. New York: Oxford University Press.
- Nahmias, Eddy, D. Coates, and T. Kvaran. 2007. "Free Will, Moral Responsibility, and Mechanism: Experiments on Folk Intuitions." *Midwest Studies in Philosophy* 31: 214–42.
- Naylor, Marjory. 1984. "Frankfurt on the Principle of Alternate Possibilities." *Philosophical Studies* 46: 249–58.
- Nelkin, Dana. 2007. "Good Luck to Libertarians." *Philosophical Explorations* 10: 173–84.
- . 2011. *Making Sense of Freedom and Responsibility*. Oxford: Oxford University Press.
- O'Connor, Timothy. 2000. *Persons and Causes*. New York: Oxford University Press.
- . 2009. "Agent-Causal Power." In T. Handfield, ed. *Dispositions and Causes*, 189–214. New York: Oxford University Press.
- . 2011. "Agent-Causal Theories of Freedom." In R. Kane, ed. *The Oxford Handbook of Free Will*, 2nd ed., 309–28. New York: Oxford University Press.
- O'Shaughnessy, Brian. 1980. *The Will*. Vol. 2. Cambridge: Cambridge University Press.
- Palmer, David. 2013. "Event-Causal Libertarianism: Two Objections Reconsidered." In I. Haji and J. Caouette, eds. *Free Will and Moral Responsibility*, 98–122. Newcastle-upon-Tyne: Cambridge Scholars Publishing.
- Pereboom, Derk. 2001. *Living without Free Will*. Cambridge: Cambridge University Press.
- . 2014. *Free Will, Agency, and Meaning in Life*. Oxford: Oxford University Press.
- . 2015. "The Phenomenology of Agency and Deterministic Agent Causation." In H. Pedersen and M. Altman, eds. *Horizons of Authenticity in Phenomenology, Existentialism, and Moral Psychology*, 277–94. Dordrecht: Springer.
- Pérez de Calleja, Mirja. 2014. "Cross-World Luck at the Time of Decision Is a Problem for Compatibilists as Well." *Philosophical Explorations* 17: 112–25.

- Pink, Thomas. 1996. *The Psychology of Freedom*. Cambridge: Cambridge University Press.
- . 2011. "Freedom and Action without Causation: Noncausal Theories of Freedom and Purposive Agency." In R. Kane, ed. *The Oxford Handbook of Free Will*. 2nd ed., 349–65. Oxford: Oxford University Press.
- Robinson, Michael. 2012. "Modified Frankfurt-Style Counterexamples and Flickers of Freedom." *Philosophical Studies* 157: 177–94.
- Rorty, Amelie. 1980. "Where Does the Akratic Break Take Place?" *Australasian Journal of Philosophy* 58: 333–46.
- Scanlon, Thomas. 1998. *What We Owe to Each Other*. Cambridge, Mass: Harvard University Press.
- Schlick, Moritz. 1962. *Problems of Ethics*. David Rynin, trans. New York: Dover.
- Searle, John. 1983. *Intentionality*. Cambridge: Cambridge University Press.
- . 2001. *Rationality in Action*. Cambridge, Mass.: MIT Press.
- Schon, Scott. 1994. "Teleology and the Nature of Mental States." *American Philosophical Quarterly* 31: 63–72.
- . 1997. "Deviant Causal Chains and the Irreducibility of Teleological Explanation." *Pacific Philosophical Quarterly*. 78: 195–213.
- . 2005. *Teleological Realism: Mind, Agency, and Explanation*. Cambridge, Mass.: MIT Press.
- Shoemaker, David. 2011. "Attributability, Answerability, and Accountability: Toward a Wider Theory of Moral Responsibility." *Ethics* 121: 602–32.
- Stapp, Henry. 2007. *Mindful Universe*. Berlin: Springer.
- Steward, Helen. 2012. *A Metaphysics for Freedom*. Oxford: Oxford University Press.
- Strawson, Galen. 2002. "The Bounds of Freedom." In R. Kane, ed. *The Oxford Handbook of Free Will*, 441–60. New York: Oxford University Press.
- Stump, Eleonore. 1990. "Intellect, Will and the Principle of Alternate Possibilities." In M. Beaty, ed. *Christian Theism and the Problems of Philosophy*, 254–85. Notre Dame, IN: University of Notre Dame Press.
- Stump, Eleonore, and N. Kretzmann. 1991. "Prophecy, Past Truth, and Eternity." *Philosophical Perspectives* 5: 395–424.
- Taylor, Richard. 1966. *Action and Purpose*. Englewood Cliffs, N.J.: Prentice-Hall.
- Thalberg, Irving. 1977. *Perception, Emotion, and Action*. New Haven, Conn.: Yale University Press.
- Thomson, Judith. 1977. *Acts and Other Events*. Ithaca, N.Y.: Cornell University Press.
- Tognazzini, Neal. 2011. "Understanding Source Incompatibilism." *Modern Schoolman* 88: 73–88.
- van Inwagen, Peter. 1983. *An Essay on Free Will*. Oxford: Clarendon Press.
- . 2000. "Free Will Remains a Mystery." *Philosophical Perspectives* 14: 1–19.
- . 2011. "A Promising Argument." In R. Kane, ed. *The Oxford Handbook of Free Will*. 2nd ed., 475–83. New York: Oxford University Press.
- Vargas, Manuel. 2012. "Why the Luck Problem Isn't." *Philosophical Issues* 22: 419–36.
- Velleman, J. David. 1992. "What Happens When Someone Acts?" *Mind* 101: 461–81.

- Wallace, R. Jay. 1999. "Three Conceptions of Rational Agency." *Ethical Theory and Moral Practice* 2: 217–42.
- Warfield, Ted. 2003. "Compatibilism and Incompatibilism: Some Arguments." In M. Loux and D. Zimmerman, eds., *The Oxford Handbook of Metaphysics*, 613–30. Oxford University Press.
- Watson, Gary. 1996. "Two Faces of Responsibility." *Philosophical Topics* 24: 227–48.
- Williams, Bernard. 1993. *Shame and Necessity*. Berkeley: University of California Press.
- Wilson, George. 1989. *The Intentionality of Human Action*. Stanford, Calif.: Stanford University Press.
- . 1997. "Reasons as Causes for Action." In G. Holmström-Hintikka and R. Tuomela, eds. *Contemporary Action Theory*, vol. 1, pp. 65–82. Dordrecht: Kluwer.
- Yates, J. Frank. 1990. *Judgment and Decision Making*. Englewood Cliffs, N.J.: Prentice-Hall, Inc.
- Zagzebski, Linda. 1991. *The Dilemma of Freedom and Foreknowledge*. Oxford: Oxford University Press.

Index

- ability, 63–86. *See also* ability to *A*
intentionally (*I*-ability); ability to
do otherwise; promise-level ability
(*P*-ability); simple ability (*S*-ability)
and ability to *A* at will, 72–73
and ability to try, 71, 174
and abnormal sources of belief about,
74, 78–80
and determinism, 65–73
and ensurance-level ability, 71–74
and freedom-level ability, 66, 87n9
general vs. specific, 63, 65
hypotheses about, 76–83
and *L*-ability, 66–69, 87n8
and meaning of “able”, 63–65, 86n5,
121, 132n6
and normal vs. abnormal
circumstances, 66–68, 75, 83
and *O*-ability, 116–17, 128, 130, 231–33
and obstacles, 80–83, 88–89nn28–31,
110, 117
reliability of, 63, 69–75, 80–82, 89n31,
110, 116–17, 127
ability to *A* intentionally (*I*-ability),
63–69, 75–79, 83–86, 87n14,
88–89n30, 117–18, 125, 133n11
ability to do otherwise, 67–69, 84–86,
87n12, 100, 116–17, 183–84, 196n13
and possible worlds, 67–68
action, 1–3. *See also* basic action;
intentional action
individuation of, 1–3, 60n18, 87n10,
234nn5 and 6
overt, 14, 56–57, 85, 111–12, 121–23, 127,
132, 234n3
and time, 86n7
Adams, F., 25n11, 59nn4 and 5,
88n19, 133n12
agent-causal libertarianism, 4, 105,
134n18, 139–40, 151n1, 155–59,
165–66, 172, 179n5, 186–87
agent causation, 4, 69, 79, 87n16, 104–6,
110, 167, 175–76, 181–82, 185–86,
195n1, 197, 210–11, 222, 228
and action, 164–65, 173
and control, 143–47, 150, 169–73,
223–25, 228, 233
and deciding, 161–69, 233
and evidence, 144–47, 185–87
and luck, 115, 118–19, 126, 129, 136,
139–40, 150–51, 155–59, 170–73, 215,
223–26, 233
nature of, 5, 179–80n5
as primitive, 118
and settling, 161–62, 178–79, 225
skepticism about, 115, 129, 132nn1 and
5, 146–49, 166, 180n9, 185–86,
192–93, 221–22, 226

- akratic action, 10, 22, 84
 Anscombe, G. E. M., 1
 Armstrong, D., 59n4
 Audi, R., 11, 24n3, 107n7
 Austin, J., 86n2
 Ayer, A., 107n7
- Balaguer, M., 167
 Baron, R., 162–63
 basic action, 11, 22, 73, 85, 133n11
 and direct control, 229–33, 234n6,
 234–35n7, 235n8
 basically free action, 112–13, 198, 203–4,
 207–11, 217, 218nn3, 6 and 8–10
 defined, 112
 and *LG*, 204
 basically free decision, 112–13, 203–4,
 207, 209, 217, 218n10
 and *LD*, 203–4, 218n6
 Berofsky, B., 135–39, 147–48
 Bishop, J., 226–29, 234n7
 Bob and the coin examples, 112–14,
 136–39, 168–72, 202–3, 206–11, 217,
 218n12, 233
 Brand, M., 11, 14, 25n13, 30
 Bratman, M., 11, 14, 25nn9, 11 and
 13, 133n10
 Brembs, B., 181–82
- C agents, 79–84, 89n31
 defined, 79–80
 Campbell, C., 94
 Cappelen, H., 24n6
 causalism, 27–58, 60n21
 and anticausalism, 27, 32–58
 and causation, 31–32
 a challenge for, 57
 and fact causation, 30–33
 and neural states, 30–31, 46–54
 and reasons, 27–31
 and states of mind, 28–33
 sufficient condition for truth of, 30–31
 causation, 32, 47, 161, 205, 225.
 See also agent causation;
 proximal causes
 and causalism, 32
 and compulsion, 100–101
 by facts, 30–33, 46, 57
 by neural states, 30–33, 37, 46–54, 57
 by reasons, 27–31, 46, 57
 by states of mind, 28–33, 37, 46–54, 57
- Child, W., 30
 Chisholm, R., 5
 Clarke, R., 5, 52–53, 94, 119–25,
 129–32, 132nn1, 5, and 8, 133nn9
 and 13, 141–49, 171–72, 185–86, 193,
 195n7, 221–29
 Coates, D., 134n17
 Cohen, J., 96
 compatibilism about free will, 4, 6n4,
 65–67, 97–107, 126–28, 149–50
 and ability to do otherwise, 67–68,
 87n12, 116–17
 agnosticism about, ix, 140, 147, 150,
 191–92, 196n6, 197
 and an analogue of *L*-ability, 67
 defined, 4
 and semicompatibilism, 68–69,
 87n13, 97–98
 compatibilism about moral
 responsibility, 68–69, 97–107
 agnosticism about, 140
 Compatibilism or Bust, 191–92
 Compatibilism Wins a Battle, 191
 complete control, 115, 154, 173–76
 and agent causation, 115, 159, 224–35
 over decisions, 158–63, 167–79, 224–25
 and direct control, 231–32
 and free action, 159–60, 175, 225
 meaning of, 115, 170
 and settling, 159, 167, 224–25
 contrastive explanation, 171–72
 control, 35–41, 64–65, 70–71, 76–85,
 88–89n30, 92, 102–6, 112–15,

- 135–42. *See also* complete control;
direct control; guidance control;
lame-control argument; more-
control argument; regulative
control; same-control argument
amounts of, 143–47, 152n5, 185,
192, 224
compatibilist, 146, 223
and determinism, 64–65
and directly free actions, 109
freedom-level, 145–46
indeterministic, 143–47, 183–85,
189–90, 195–96n9, 196n10
moral-responsibility-level, 137–39
counterparts, 68, 87n11, 156, 179n4
- Dancy, J., 27–28
- daring libertarian view (*DLV*), 197, 199,
205–17, 218n10, 225–26, 234
and agents' histories, 210–17
and daring soft libertarianism,
217–18n1
vs. modest libertarian view, 197
- Davidson, D., 1–3, 6n2, 14, 24n3, 27–33,
37–38, 43–47, 54–56, 235n7
- decide signal, 174–76
- deciding, 7–23, 24n5.
See also decision; intention
arbitrary, 94, 136–39
experience of, 12–14
as extended action, 8–11
intentions in, 14–18
as mental action, 7–23
as mythical, 9–13
and neuroscience, 26nn21 and 23
as nonactional, 8–13
practical vs. cognitive, 7–13, 20–23
questions about, 7
scope of, 23
and settling, 10–22, 24–25n9, 88n18,
111, 128, 130, 155, 165–67
special nature of, 123–24
and uncertainty, 12–13, 18–22, 24n7,
26nn19 and 22, 92
- decision, 3, 7. *See also* deciding
and decision not to *A*, 7, 18, 21–22
and decision state, 155, 165–67, 179n2,
179–80n5
proximal, 137–38, 174
and referents of “decision”, 155
- Dennett, D., 182
- determinism, 4, 97–101, 106, 112, 115–18,
126, 134n17, 165, 181–94, 198, 230
and ability, 65–69, 116
and control, 64–65, 134n18,
141–42, 146–50
defined, 4
and sufficient conditions for free
action, 140–41, 147–49
- direct control, 114–15, 142–44, 157,
226–33, 234nn2–5, 234–35n7, 235n8
and agent-causation, 228, 233
and basic action, 229–33, 234n6,
234–35n7, 235n8
- Bishop-style and Knight-style
conceptions of, 228–33, 235n10
and complete control, 231–32
and deciding, 228–33, 235n10
and free action, 229–33, 235n8
and indirect control, 152n5, 161–62,
227, 230–32
and luck, 231–33
supposition *S* about, 229–30, 235n8
- directly free action, 4, 65, 86n6, 92–93,
99–103, 109–15, 120–32, 133n13,
136–37, 151, 183–94, 205, 218nn6 and
8, 221–26, 235n8
and decisions, 121
and *LDF*, 109, 115, 136, 151, 168–76, 186
and *LDFd*, 123, 129–30, 218n6, 221, 226
- disappearing agent objection, 153–67,
178–79, 179n1, 186, 221,
224–26. *See also* settling whether a
decision occurs

- disappearing agent objection (*Cont.*)
 and agency, 163–66
 and basic desert moral responsibility,
 153–55, 158, 161–67, 177–79, 224
 and complete control, 158–79,
 224–25
 formulation of, 153
 and luck, 158–59, 178, 226
 modified formulation of, 166
 D'Oro, G., 45–46, 58
- effort, 35, 88n18, 201
 dual or concurrent, 197–210, 217,
 218nn5 and 12
 and trying, 35–37, 50, 92
- event-causal libertarianism, 4–6, 106,
 114–15, 118, 126, 140, 197, 200, 205,
 210, 217, 234. *See also* disappearing
 agent objection; lame-control
 argument; more-control argument;
 same-control argument
 and control, 42, 104–5, 141–51,
 185–86, 221–27, 233
 and *LFTe*, 221, 234n1
- Fischer, J., 68–69, 86n4, 94, 97–98,
 107n9, 142, 152n7
- Frankfurt, H., 23, 24n1, 183, 195n3
- Frankfurt-style stories, 98–102, 106,
 107n6, 127–29, 183–84, 187–90,
 193–94, 195nn4 and 8, 195–96n9,
 196n13, 218n10
 global, 102, 107n9, 189
 free action, 4. *See also* basically free
 action; directly free action;
 indirect, 6n5, 65, 86n6,
 92–93, 133n9, 168
 and leeway incompatibilism,
 101–3, 106
 and moral-responsibility-level free
 action, 4, 6n3, 106,
 135–39, 191
- freedom to do otherwise, 97–102, 107n5
- free will, 5. *See also* basically free
 action; directly free action; free
 action; moral responsibility, and
 connection to free will
 importance of, 103–5
 meaning of, 4, 92, 96–97
 and mystery, 84, 110, 126
 setting the bar for, 148–49, 159, 178,
 193, 221–22
- fruit flies, 182–85
- Ginet, C., 2, 43, 47–54, 58, 59nn10–11, 13,
 and 14, 131
- Goldman, A., 1
- Greene, J., 96
- Griffith, M., 151, 167–69, 177, 224
- Grünbaum, A., 107n7
- guidance control, 98, 142, 146, 149,
 183–85, 192, 221–24
- Haji, I., 155, 167
- Hornsby, J., 2, 30, 88n19, 234n3
- Hume, D., 107n7
- incompatibilism, 4, 6n4, 99–106, 132n5,
 147–49, 172, 182–97, 205, 213,
 233–34. *See also* source
 incompatibilism
- indeterministic process, 182–83, 186–87,
 191–92, 213
 early, 182–83, 195n5
 evidence of, 181–82, 186, 191, 205, 234
 and evolution, 182–83, 186
 late, 182–84, 187, 193–94
- intention, 3
 belief constraints on, 20, 75
 content of, 18–22, 26n18, 47,
 54–57, 60n15
 about decision, 14–18, 24n8, 71–74,
 127, 160–61
 by default, 10
 and desire, 19
 distal, 3, 16

- executive aspect of, 18–20
- and guidance, 19, 25n14
- mixed, 3
- nonactionally acquired, 8–17, 20–23, 26n21, 121, 195n5
- orientation of, 19–20
- and plans, 18–21, 25n13
- proximal, 3, 11–12, 16, 26n21, 73, 232
- and settledness, 12–16, 19–21, 24–25n9, 25n15
- intentional action, 2–3, 6nn1 and 2, 14, 17, 28–32, 49, 55, 60n21, 64, 71, 84–86, 116, 133n11, 173, 230
- and essentially intentional action, 84, 89n34, 209
- explanations of, 27–58, 58n1
- and intention, 14, 17, 19, 23, 25n11
- intuition, 24n6
- Jahanshahi, M., 26n21
- James, W., 59n5
- Kane, R., 5, 17, 24n1, 94, 107n2, 130–31, 143–44, 157, 169, 181, 192, 197–210, 215–17, 218nn2, 4–5, 8–9, 11, and 12
- Kapitan, T., 86n3
- Kaufman, A., 7, 24n1
- Kavka, G., 25n10, 26n19
- Kearns, S., 132n2, 152n7
- Kretzmann, N., 183
- Kvaran, T., 134n17
- lame-control argument, 145–46
- Lau, H., 26n21
- LDF*, 109, 115, 136, 151, 169–70, 173, 176, 186
- LDFd*, 123, 129–30, 218n6, 221, 226
- Levy, N., 152n3
- Lewis, D., 31–32
- LF*, 135, 149–50, 190–94
- LFT*, 135, 147, 150, 151n1, 181, 186, 190–94, 234n1
- LFTe*, 221, 234n1
- libertarianism, 4–6, 84–85, 96, 107–15, 127–29, 134nn16 and 18. *See also* agent-causal libertarianism; agent causation; daring libertarian view; event-causal libertarianism; noncausal libertarianism; soft libertarianism
- and ability to do otherwise, 65–66, 87n12
- defined, 4
- and directly vs. indirectly free action, 65, 86n6
- and ensurance-level ability, 71–74
- kinds of, 4
- and *L*-ability, 66–69, 87n8
- and *LDF*, 109, 115, 136, 151, 169–70, 173, 176, 186
- and *LDFd*, 123, 129–30, 218n6, 221, 226
- and *LF*, 135, 149–50, 190–94
- and *LFT*, 135, 147, 150, 151n1, 181, 186, 190–94, 234n1
- and mystery, 84, 126
- and *P*-ability, 74–86, 110, 116–18, 127, 132
- and sufficient conditions for free action, 141
- luck, 109, 112. *See also* problem of present luck
- in deterministic scenarios, 151–52n3
- and intentional action, 64, 85–86
- just a matter of, 109, 115, 119–32, 136, 151, 156–57, 163, 169–70, 176–77, 186, 198–99, 203–4, 208, 221–26, 231
- meaning of, 139–40, 167–70
- partly a matter of, 115, 118–19, 151, 156–59, 168–69, 176–77, 198
- recipient of, 135–40
- Mavrodes, G., 94
- Maye, A., 181–82
- McCann, H., 14–16, 24nn1, 8, and 9, 25nn11 and 17, 54, 59n4, 209

- McGinn, C., 59n4
 McKenna, M., 107n8, 195n5
 Mellor, D. H., 30
 Mickelson, K., 6n4
 Mill, J. S., 107n7
 Miller, J., 24n5
 Montague, P. R., 96, 148
 moral responsibility, 4, 6n3, 68–69, 84,
 91–107, 137–40, 170–71, 184,
 188–91, 194–95, 195n4, 196n13
 accountability sense of, 91
 basic, 113, 212–17, 219n16
 and basic desert, 153–67, 177–79, 234
 of children, 212–17
 and connection to free will, 91–107
 and contrastive facts, 204–5
 direct, 93, 100–103, 113–14, 136, 171–72
 indirect, 93–94
 leeway incompatibilism about, 103
 and source incompatibilism, 101–2,
 106, 107n8
 sphere of, 94–95
 more-control argument, 143–46,
 185–86, 221–24
 Moser, P., 60n21, 64, 85

 Nagel, T., 164
 Nahmias, E., 134n17
 Naylor, M., 196n13
 Nelkin, D., 195n1
 noncausal libertarianism, 4, 6n5,
 172–73, 211
 not-doing, 21, 23, 88n17
 and deciding not to *A*, 7–13, 18, 21–23
 and intention not to *A*, 21, 24n9, 25n15

 O'Connor, T., 5, 6n6, 94, 114–15, 118–19,
 126, 129, 143, 152n6, 156–57, 170–73,
 180n5, 226
 O'Shaughnessy, B., 12, 24n2

 Pereboom, D., 5, 126, 132n1, 141–49,
 153–67, 177–79, 179nn1 and 5,
 180nn7 and 9, 183–86, 190–92,
 195nn1 and 7, 221–26
 Pérez de Calleja, M., 152n3
 Pink, T., 6n6, 24n1
 practical reflection, 8–13
 problem of present luck, 112–15,
 118–31, 133nn9 and 13, 134n16,
 135–51, 151n3, 152n7, 191, 203–4,
 218n12, 226
 and agent causation, 118–19, 126, 129,
 136, 139–40, 143–47, 150–51, 155–59,
 161–73, 215, 223–26, 233
 and arbitrary decisions, 136–39
 and complete control, 115, 167–79
 and continuations, 113–14
 and problem of evil, 136, 171, 203–4
 solution to, 115, 129, 136–38, 141,
 148–51, 156, 171, 176–78, 198–200,
 210–17, 233
 statement of, 112–14
 promise-level ability (*P*-ability), 74–86,
 88–89nn30–32, 110, 116–18, 127, 132
 promising, 70, 74–83, 88n24,
 89n32, 127
 and confidence, 74–75, 78–82, 85,
 88–89n30, 89n31, 110, 116–17, 127
 and control, 76–78
 and deciding, 74, 127
 and excuses, 79–80, 83, 89n32
 and intention, 70, 74–77
 and manipulation, 79–80
 proximal causes, 4, 65, 115, 135, 140–41,
 150, 151n1, 181, 184–90, 194,
 198–200, 205, 218n8, 219n16, 221

 randomizer, 73, 85, 113, 160, 174, 211,
 216, 231–32
 Ravizza, M., 69, 94, 98
 reasons, 27–58, 60n18
 for deciding, 21, 25n17, 26n18, 55–56,
 60nn19 and 20
 objective favorers view of, 28–29, 58n1
 prevailing of, 199–203, 207–9, 218n12

- regulative control, 142, 146, 149–50, 183,
192, 221–26
- Robb, D., 98, 102, 107n9,
195n4
- Robinson, M., 196n13
- Rorty, A., 24n3
- same-control argument, 142–43, 146–47,
186, 221–24
- Scanlon, T., 27–28
- Schlick, M., 107n7
- Schwarz, W., 24n5
- Searle, J., 19, 24n1
- Schon, S., 31, 37–45, 58
- semicompatibilism, 68–69, 87n13,
97–98, 106
- sentient direction, 34–37, 41, 51
- settling, 10, 13, 18, 20, 153–55, 166.
See also settling whether a
decision occurs
in deciding, 20–22, 155,
165–67
and decision states, 155, 165–67,
179–80n5
and settledness, 10–22, 24–25n9,
25nn15 and 16, 166
- settling whether a decision occurs,
153–55, 158–67, 177–79, 179–80n5,
180n7, 224–25
and complete control, 158–70,
174–79, 224–25
- Shoemaker, D., 91
- simple ability (*S*-ability), 63–74,
78, 84–86
- soft libertarianism, 194–95,
217–18n1, 235n11
- source incompatibilism, 101–2, 106,
107n8, 195nn5 and 6
- Stapp, H., 183
- Steward, H., 115
- Strawson, G., 214
- Stump, E., 183
- Taylor, R., 5, 31, 173
- teleological explanations, 31,
37–44, 57–58
- Thalberg, I., 2
- Thomson, J., 2
- Tognazzini, N., 107n8
- toxin puzzle, 25n10, 26n19
- trying, 34–35, 71–73, 92, 133n12
- up to an agent, 109, 112–32, 134n17, 175,
189, 201–9, 231.
See also *UT*; *UTD*; *UTn*
and *LDF*, 109
and *LDFd*, 129–30
and luck, 113–32
UT, 120–22, 133n11, 133–34n14
UTD, 122–26, 133–34n14
UTn, 129–31
- van Inwagen, P., 4, 69–71, 74, 84–86,
87n16, 89n34, 97, 100, 107n4,
110–12, 115–18, 126–29, 132nn1 and
6, 187, 192, 196n13, 227
- Vargas, M., 152n7
- Velleman, J. D., 164
- volition, 49–52, 59n14
- voting examples, 85–86, 91–93, 133n9,
168–69, 175
- Wallace, R. J., 43, 47, 54–58,
60nn15 and 19
- Warfield, T., 98
- Watson, G., 91
- Williams, B., 24n2
- willing, 49, 72–73. *See also* volition
- Wilson, G., 31–45, 51–52, 58
- Yates, J. F., 7
- ZAF*, 188–89
- Zagzebski, L., 183
- zygote argument. *See* *ZAF*

